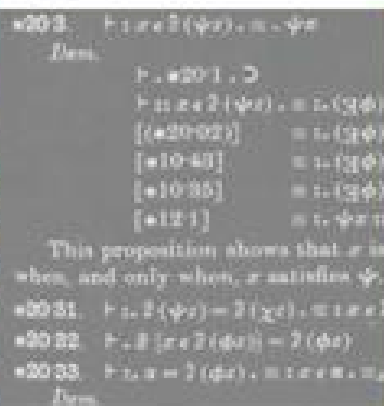
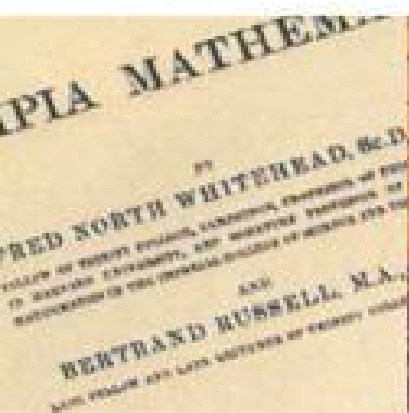


BERTRAND RUSSSELL

INTRODUÇÃO À FILOSOFIA MATEMÁTICA





Bertrand Russell

Introdução à filosofia da matemática

Tradução:
Maria Luiza X. de A. Borges

Revisão técnica:
Samuel Jurkievicz
UFRJ



Sumário

Introdução, de John G. Slater

Prefácio

1 • A série dos números naturais

2 • Definição de número

3 • Finitude e indução matemática

4 • A definição de ordem

5 • Tipos de relações

6 • Similaridade das relações

7 • Números racionais, reais e complexos

8 • Números cardinais infinitos

9 • Séries e ordinais infinitos

10 • Limites e continuidade

11 • Limites e continuidade das funções

12 • Seleções e axioma multiplicativo

13 • O axioma da infinidade e tipos lógicos

14 • Incompatibilidade e a teoria da dedução

15 • Funções proposicionais

16 • Descrições

17 • Classes

18 • Matemática e lógica

Índice remissivo

Introdução

Russell escreveu *Introdução à filosofia da matemática* na prisão, durante o verão de 1918. Em janeiro daquele ano ele havia sido intimado pelas autoridades e acusado de fazer uma declaração que insultava um aliado da Grã-Bretanha em tempo de guerra. O aliado eram os Estados Unidos, e o que Russell escrevera era que os Estados Unidos provavelmente tenderiam a intimidar grevistas na Grã-Bretanha e na França depois que a guerra terminasse, “uma ocupação que o Exército norte-americano está acostumado a ter em casa”. Embora ele baseasse sua declaração num relatório do Senado dos Estados Unidos, foi julgado, condenado e sentenciado a seis meses como prisioneiro na segunda divisão. (Talvez fossem as duas frases seguintes à citada que tivessem levado as autoridades, que haviam suportado provocações de Russell durante dois anos, a processá-lo: “Não digo que estes pensamentos estejam na mente do governo. Todos os indícios tendem a mostrar que não há absolutamente nenhum pensamento na mente daqueles que ocupam os postos do Estado, e que eles vivem ao deus-dará, consolando-se com a ignorância e as tolices sentimentais.”) A perspectiva de passar seis meses na segunda divisão deixou Russell e seus amigos extremamente alarmados, temendo danos para seu intelecto; assim, pressionaram o governo para que a sentença fosse alterada pelo encarceramento na primeira divisão.

Naquele tempo as distinções de classe penetravam até as prisões. Os prisioneiros da primeira divisão pagavam aluguel pelo uso de suas celas; tinham permissão para mobiliá-las com seus próprios móveis; estavam autorizados a contratar outro prisioneiro para lhes servir de criado; não eram submetidos ao racionamento de comida; tinham permissão para dispor de livros e material para escrever. Todos esses privilégios eram negados aos confinados na segunda divisão. Além disso, os prisioneiros da primeira divisão tinham direito a receber mais cartas e mais visitas que os que cumpriam pena na segunda. Para seu mérito perene, Arthur Balfour, cujas pretensões e capacidades filosóficas Russell e difamara, e a cujas

políticas públicas opunha-se acerbamente, conseguiu que a sentença fosse alterada. Enormemente aliviado, Russell passou a planejar o melhor uso possível para sua retirada forçada do mundo.

Havia alguns anos que ele pretendia escrever um livro-texto de lógica. Tinha aguda consciência de que *Principia Mathematica*, apesar de sua reconhecida importância, tinha muito poucos leitores; estava convencido, no entanto, de que, se novos filósofos o compreendessem, enfrentariam os problemas filosóficos de maneira muito mais profícua do que costumavam fazer. O mais importante nos *Principia* para efetuar essa revolução são seus conceitos básicos que, Russell estava convencido, podem ser apreendidos independentemente da massa de símbolos em que estão inseridas naquele livro. Na iminência da prisão, ele decidiu que chegara o momento de levar a cabo seu projeto. O plano amadurecido requeria dois livros: o primeiro, um “Prelúdio aos *Principia*”; o segundo, uma meticulosa reelaboração de “A filosofia do atomismo lógico”. Como a primeira parte do projeto era simplesmente uma introdução aos *Principia*, nenhuma pesquisa adicional se faria necessária. Ele precisaria apenas organizar o material que tinha na cabeça e pô-lo no papel.

Nos meses que precederam imediatamente sua intimação, Russell havia proferido duas séries de conferências em Londres para públicos pagantes. A primeira delas abrangeu o mesmo terreno que *Introdução à filosofia da matemática*; a segunda tornou-se famosa como “A filosofia do atomismo lógico”. Não existe nenhum manuscrito ou texto datilografado da primeira série de conferências; um estenógrafo registrou a segunda, tal como pronunciada, incluindo as discussões que se seguiram, e esse texto datilografado foi editado por Russell e outros para publicação numa revista. O primeiro conjunto, que tratava de assuntos muito conhecidos, provavelmente não diferia em substância do texto deste livro. Depois que Russell se decidia por uma maneira de explicar uma idéia, invariavelmente perseverava nela.

Em 1º de maio de 1918, o apelo de Russell contra a sentença de prisão foi rejeitado, e nesse mesmo dia ele ingressou na prisão Brixton. Ficou decepcionado quando as autoridades lhe ordenaram pegar um táxi; alimentara a esperança de que providenciariam seu transporte num camburão. Em sua *Autobiografia*, rememorou a recepção que teve no portão da penitenciária:

Fui muito confortado, ao chegar, pelo carcereiro no portão, que havia obtido informações sobre mim. Perguntou minha religião e respondi “agnóstico”. Pediu-me que soletrasse a palavra e comentou, com um suspiro: “Bem, há muitas religiões, mas suponho que todas cultuam o mesmo Deus.” Esse comentário me manteve de bom humor por cerca de uma semana.

Antes de ser encarcerado, ele havia traçado seu plano de trabalho: quatro horas de escrita filosófica por dia, quatro horas de leitura filosófica e quatro horas de leituras gerais. Na segunda-feira, 6 de maio, porém, ainda não tinha livros nem material para escrever. Havia, contudo, mobiliado sua cela com uma cama e outros móveis enviados pelo irmão, a quem escreveu nesse dia: “Espero logo ter material para escrever: então redigirei primeiro um livro chamado *Introdução à lógica moderna*, e quando ele estiver terminado começarei uma obra ambiciosa que deverá se chamar *Análise da mente*. As condições aqui são boas para a filosofia.” Mais adiante, na mesma carta, pediu ao irmão para transmitir esse recado: “Diga a Whitehead que quero escrever um livro-texto para os *Principia* e lerei tudo o que ele considerar relevante.”

Em 21 de maio Russell enviou uma mensagem a H. Wildon Carr, na época secretário honorário da Aristotelian Society, que estava atuando temporariamente como seu agente literário: “Escrevi cerca de 20 mil palavras da *Introdução à filosofia da matemática*, para desenvolver as conferências dadas *antes* do Natal. Depois trabalharei sobre as conferências dadas *depois* do Natal (que tenho, obrigado).” Mais adiante, na carta, retornou ao livro: “Espero terminar a *Introdução* dentro de mais ou menos um mês. A prisão é um bom lugar para leituras e trabalho fácil, mas seria impossível para um pensamento realmente difícil.” Apenas seis dias depois, numa carta ao irmão, Russell incluiu outro recado para Carr: “Quase terminei a *Introdução à filosofia da matemática*, um prelúdio de 70 mil palavras aos *Principia*.” Passou então a mencionar várias questões filosóficas para as quais precisava encontrar respostas defensáveis antes de se achar em condições de escrever o segundo livro projetado. Este, que se chamaria *Elementos de lógica*, “proporá a base lógica do que chamo de ‘atomismo lógico’ e situará a lógica em relação à psicologia, matemática etc.”. Pretendia tratar de muitas das questões filosóficas que o estavam perturbando em um outro projeto de livro, *Análise da mente*, e o trabalho

preliminar sobre ele consumiu o restante de seu tempo na prisão. Esse livro foi publicado em 1921.

Numa carta anterior ao irmão, em 16 de maio, Russell descreveu *Elementos de lógica* como essencial para a compreensão das idéias a serem desenvolvidas em *Análise da mente*:

Continuarei com minha *Análise da mente*, que, caso bem-sucedida, deveria ser mais uma obra extensa e importante. Ela deverá se complementar com um livro sobre lógica: não o que estou fazendo agora, que deve ser um livro-texto, mas outro, na linha das conferências que dei depois do Natal: sem tal complemento, seria praticamente ininteligível. Prevejo pelo menos três anos de trabalho sobre esse tema.

É uma pena que ele não tenha reelaborado “A filosofia do atomismo lógico” em *Elementos de lógica*; essas conferências foram citadas com muita freqüência nos anos entre as guerras, apesar de disponíveis apenas numa revista; parece provável que teriam sido ainda mais amplamente estudadas, e assim mais influentes, se fossem mais facilmente disponíveis, sob forma revista e ampliada, num livro. “A filosofia do atomismo lógico” foi reproduzida no volume 8 de *The Collected Papers of Bertrand Russell*.

Em suas discussões da *Introdução à filosofia da matemática*, Russell descreve a obra por vezes como uma introdução aos *Principia*, como foi mencionado, e por vezes como “uma versão semipopular de *Os princípios da matemática*”, como o faz em sua *Autobiografia*. Talvez esta última descrição seja um mero lapso. *Princípios*, em retrospecto histórico, foi a primeira afirmação que Russell fez da tese de que grande parte da matemática é um ramo da lógica, tese que ele e Alfred Whitehead desenvolveram elaboradamente nos *Principia Mathematica*. (*Princípios*, porém, é muito mais que isso, pois contém extensas e importantes discussões da maior parte das questões metafísicas tradicionais.) Os três livros, portanto, são obviamente relacionados, mas *Introdução à filosofia da matemática* está tão completamente imbuído das idéias centrais dos *Principia* que parece quase descabido falar dele como mais estreitamente relacionada aos *Princípios*. Para citar apenas um exemplo: a teoria das descrições de Russell, que só foi publicada dois anos depois do lançamento de *Princípios*, é objeto de um capítulo inteiro na *Introdução*.

O manuscrito da *Introdução à filosofia da matemática* encontra-se agora nos Russell Archives, na McMaster University, em Hamilton, em Ontário. O diretor da prisão foi obrigado a lê-lo para verificar se continha alguma coisa proibida pelos regulamentos do tempo de guerra, mas, supomos que com grande alívio, delegou a tarefa Carr. Este lhe assegurou que a obra era o que seu título indicava e que nada continha de subversivo. Uma vez fora da prisão, o manuscrito foi entregue à datilógrafa costumeira de Russell, a senhorita Kyle, para que preparasse uma cópia para os tipógrafos. Numa carta de 29 de julho ao irmão, Russell expressou certa irritação com o atraso da datilografia: “Diga à senhorita Kyle para se apressar com a *Introdução à filosofia da matemática* — ela está com o livro há *realmente* muito tempo.” As cartas que enviou da prisão não registram a data em que apresentou o livro ao diretor para exame, mas presumivelmente foi no início de junho, pois o livro tem menos de 80 mil palavras. O manuscrito foi finalmente devolvido a Russell e ele o deu de presente a lady Constance Malleson, com quem mantinha um caso na época. Antes de sua morte, em 1975, ela o vendeu ao Russell Archives com todos os outros materiais de Russell que possuía. Uma inspeção desse manuscrito mostra claramente que Russell seguiu seu próprio conselho de fazer todas as mudanças em seu texto na própria cabeça e depois simplesmente redigir a forma final. Há muito poucas alterações em seu manuscrito e absolutamente nenhuma tentativa de formulação foi cancelada.

Num ponto do livro, Russell alude às condições adversas sob as quais ele foi escrito. No início do Capítulo 16, adverte o leitor de que o artigo definido exige dois capítulos: um para o significado do singular *o* e outro para o do plural *os*.

Pode ser considerado excessivo dedicar dois capítulos a uma palavra, mas para o matemático filosófico essa é uma palavra de grande importância: como o gramático de Browning¹ com o enclítico $\delta\epsilon$, eu enunciaria a doutrina dessa palavra se estivesse “morto da cintura para baixo”, e não meramente numa prisão.

Carr talvez tenha desviado os olhos ao topar com essa passagem.

A melhor descrição dos conteúdos deste livro e de seu nível de dificuldade foi escrita pelo próprio Russell como publicidade para a primeira edição. O texto apareceu apenas na sobrecapa da primeira

impressão da primeira edição; já na segunda impressão, foi substituído por um trecho de uma crítica entusiástica publicada no *Athenaeum*.

Este livro destina-se aos que não têm conhecimento prévio dos tópicos de que trata, e não sabem da matemática nada além do que pode ser adquirido na escola primária ou mesmo em Eton. Ele expõe de forma elementar a definição lógica de número, a análise da noção de ordem, a doutrina moderna do infinito e a teoria das descrições e classes como ficções simbólicas. Os aspectos mais controversos e incertos da matéria são subordinados àqueles que podem agora ser considerados conhecimento científico adquirido. Esses são explicados sem o uso de símbolos, mas de maneira a dar aos leitores uma compreensão geral dos métodos e objetivos da lógica matemática, que, como se espera, será do interesse não só dos que desejam avançar para um estudo mais sério da matéria, como daquele círculo mais amplo que sente um desejo de conhecer os procedimentos dessa importante ciência moderna.

A maliciosa referência de Russell a Eton provavelmente deixou Stanley Unwin constrangido, e, na primeira oportunidade, ele parou de usar essa descrição. Mas ela é tão quintessencialmente russelliana que merece ser muito mais conhecida do que foi até agora. Exceto por sua avaliação do preparo exigido do leitor, que certamente se pode questionar, o restante do texto descreve realmente com precisão do que trata o livro. Mas este é um livro para ser estudado, não meramente lido. Os que o estudarem cuidadosamente emergirão com uma boa apreensão da filosofia matemática de Russell e terão se preparado para enfrentar seus livros e ensaios mais técnicos, em que suas contribuições originais para a lógica matemática tinham sido publicadas anteriormente. Somente assim poderão experimentar o gênio de Russell para o trabalho original (em contraposição à mera exposição) em primeira mão.

John G. Slater
Universidade de Toronto

Prefácio

Este livro destina-se a ser essencialmente uma “Introdução”, e não pretende apresentar uma discussão exaustiva dos problemas de que trata. Pareceu desejável expressar certos resultados, até agora só disponíveis para aqueles que dominaram o simbolismo lógico, de uma forma que ofereça o mínimo de dificuldade para o iniciante. Foi feito o máximo esforço para evitar o dogmatismo acerca de questões que ainda estão abertas a séria dúvida, e esse esforço dominou em certa medida a escolha dos tópicos considerados.

Os primórdios da lógica matemática são menos explicitamente conhecidos que suas parte mais recente, mas são de interesse filosófico pelo menos igual. Muito do que é exposto nos capítulos que se seguem não deve ser propriamente chamado de “filosofia”, embora as matérias envolvidas fossem incluídas na filosofia quando não existia para elas nenhuma ciência satisfatória. A natureza da infinidade e da continuidade, por exemplo, pertenceu antigamente à filosofia, mas hoje faz parte da matemática.

Talvez não se possa argumentar que a *filosofia* matemática, no sentido estrito, inclui resultados científicos definidos como os que foram obtidos nessa região; deve-se esperar, naturalmente, que a filosofia da matemática trate de questões na fronteira do conhecimento, com relação às quais ainda não se obteve certeza relativa. Mas a especulação acerca dessas questões dificilmente tenderá a ser frutífera, a menos que as partes mais científicas dos princípios da matemática sejam conhecidas. Um livro que trate dessas partes pode, portanto, pretender ser uma *introdução* à filosofia da matemática, embora dificilmente possa almejar, exceto quando pisa fora de sua província, estar realmente lidando com uma parte da filosofia. Ele trata efetivamente, no entanto, de um corpo de conhecimento que, para aqueles que o aceitam, parece invalidar grande parte da filosofia tradicional, e até boa parte do que é corrente nos dias de hoje. Dessa

maneira, assim como por se relacionar com problemas ainda não resolvidos, a lógica matemática é relevante para a filosofia.

Por essa razão, bem como em virtude da importância intrínseca da matéria, pode ser de alguma utilidade fazer um relato sucinto dos principais resultados da lógica matemática de uma forma que não exija qualquer conhecimento de matemática, nem aptidão para simbolismo matemático. Aqui, contudo, como em outros lugares, o método é mais importante que os resultados do ponto de vista de pesquisas adicionais; e o método não pode ser bem explicado na estrutura de um livro como o que se segue. É de esperar que alguns leitores possam ficar suficientemente interessados para avançar em direção a um estudo do método pelo qual a lógica matemática pode se tornar útil na investigação dos problemas tradicionais da filosofia. Mas esse é um tema de que as páginas que se seguem não tentaram tratar.

Bertrand Russell

Capítulo 1

A série dos números naturais

A matemática é um estudo que, quando partimos de suas partes mais conhecidas, pode ser continuado em uma de duas direções opostas. A direção mais conhecida é construtiva, rumo a uma complexidade gradualmente crescente: de números inteiros para frações, números reais, números complexos; de adição e multiplicação para diferenciação e integração, e adiante, para a matemática superior. A outra direção, menos conhecida, procede, por análise, rumo à abstração e à simplicidade lógica cada vez maiores; em vez de perguntar o que pode ser definido e deduzido do que é inicialmente suposto, perguntamos que idéias e princípios mais gerais podem ser encontrados, em termos dos quais o que foi nosso ponto de partida pode ser definido ou deduzido. É o fato de seguir essa direção oposta que caracteriza a filosofia matemática em contraposição à matemática comum. Mas é preciso compreender que a distinção se dá não na matéria, mas no estado de espírito do investigador. Os antigos geômetras gregos, ao passar das regras empíricas da topografia egípcia para as proposições gerais pelas quais essas regras puderam ser consideradas justificáveis, e delas para os axiomas e postulados de Euclides, estavam fazendo filosofia matemática, segundo a definição mencionada; mas depois que os axiomas e postulados haviam sido alcançados, o uso dedutivo deles, como o encontramos em Euclides, pertencia à matemática no sentido comum. A distinção entre matemática e filosofia matemática depende, pois, do interesse que inspira a pesquisa e do estágio que a pesquisa alcançou; não das proposições a que a pesquisa diz respeito.

Podemos expressar a mesma distinção de outra maneira. As coisas mais óbvias e fáceis em matemática não são as que vêm logicamente no começo; são aquelas que, do ponto de vista da dedução lógica, começam em algum lugar no meio. Assim como os corpos mais fáceis de se ver são

os que não estão nem muito perto nem muito longe, não são nem muito pequenos nem muito grandes, assim também as concepções mais fáceis de apreender são aquelas que não são nem muito complexas nem muito simples (usando “simples” num sentido *lógico*). E assim como precisamos de dois tipos de instrumento, o telescópio e o microscópio, para a ampliação de nossas capacidades visuais, também precisamos de dois tipos de instrumento para a ampliação de nossas capacidades lógicas, um para nos fazer avançar até a matemática superior, o outro para nos levar de volta aos fundamentos matemáticos das coisas que tendemos a dar como certas em matemática. Vamos descobrir que, ao analisar nossas noções matemáticas comuns, adquirimos nova penetração, novas capacidades e meios para alcançar temas matemáticos inteiramente novos, adotando novas linhas de avanço após nossa viagem para trás. O objetivo deste livro é explicar a filosofia matemática de uma maneira simples e não técnica, sem se aprofundar sobre aquelas partes tão duvidosas ou difíceis que se torna praticamente impossível submetê-las a um tratamento elementar. Um tratamento completo será encontrado em *Principia Mathematica*; ¹ o tratamento exposto no presente volume destina-se a ser meramente uma introdução.

Para a pessoa medianamente instruída de nossos dias, o ponto de partida óbvio da matemática seria a série dos números inteiros,

$$1, 2, 3, 4, \dots \text{etc.}$$

Provavelmente só alguém com algum conhecimento matemático pensaria em começar com 0 em vez de 1, mas vamos presumir esse grau de conhecimento; tomaremos como nosso ponto de partida a série

$$0, 1, 2, 3, \dots n, n+1, \dots$$

e é essa série que teremos em mente quando falarmos da “série dos números naturais”.

Só num estágio elevado da civilização teríamos podido tomar essa série como nosso ponto de partida. Devem ter sido necessárias muitas eras para se descobrir que uma parêntese e um par de dias eram casos do número 2: o grau de abstração envolvido está longe de ser fácil. A descoberta de que 1 é um número deve ter sido difícil. Quanto ao 0, trata-se de uma adição muito recente; os gregos e os romanos não possuíam

esse dígito. Se estivéssemos nos aventurando na filosofia matemática em tempos mais antigos, teríamos de começar com algo menos abstrato que a série dos números naturais, a qual teríamos de alcançar como um estágio em nossa viagem ao passado. Quando os fundamentos da matemática tiverem se tornado mais conhecidos, seremos capazes de começar mais atrás, no que é agora um estágio tardio de nossa análise. No momento, porém, os números naturais parecem representar o que é mais fácil e mais conhecido em matemática.

Embora conhecidos, eles não são contudo compreendidos. Muito pouca gente tem condições de dar uma definição do que se entende por “número”, ou “0”, ou “1”. Não é muito difícil ver que, começando-se por 0, se pode chegar a qualquer outro dos números naturais por adições repetidas de 1, mas precisamos definir o que entendemos por “adicionar 1” e o que entendemos por “repetidas”. Essas questões nada têm de fáceis. Até recentemente, acreditava-se que pelo menos algumas dessas primeiras noções de aritmética deviam ser aceitas como simples e primitivas demais para serem definidas. Como todos os termos que são definidos o são por meio de outros termos, é claro que o conhecimento humano deve sempre se contentar em aceitar alguns termos como inteligíveis sem definição, de maneira a ter um ponto de partida para suas definições. Não é claro que deva haver termos que sejamos *incapazes* de definir; é possível que, por mais que recuemos em nossas definições, *possamos* sempre ir mais longe. Por outro lado, é também possível que, quando a análise foi levada suficientemente longe, possamos alcançar termos que são realmente simples, e portanto não passíveis logicamente do tipo de definição que consiste em análise. Essa é uma questão que não temos necessidade de decidir; para nossos objetivos, é suficiente observar que, como as capacidades humanas são finitas, as definições conhecidas por nós sempre começam em certo ponto, com termos indefinidos no momento, embora talvez não permanentemente.

Toda a matemática pura tradicional, incluindo a geometria analítica, pode ser encarada como consistindo inteiramente em proposições acerca dos números naturais. Isto é, os termos que ocorrem podem ser definidos por meio dos números naturais, e as proposições podem ser deduzidas das propriedades dos números naturais — com a adição, em cada caso, das idéias e proposições da lógica pura.

É uma descoberta razoavelmente recente que toda matemática pura tradicional pode ser derivada dos números naturais, embora se suspeitasse disso havia muito. Pitágoras, que acreditava que não só a matemática, mas todas as outras coisas, podiam ser deduzidas dos números, foi quem descobriu o mais sério obstáculo no caminho da chamada “arimetização” da matemática. Foi ele que descobriu a existência de incomensuráveis, e, em particular, da incomensurabilidade do lado de um quadrado e sua diagonal. Se o comprimento do lado é 1 metro, o número de metros na diagonal é a raiz quadrada de 2, que não era número algum. O problema então suscitado só foi resolvido em nosso próprio tempo, e só foi resolvido *completamente* com a ajuda da redução da aritmética à lógica que será explicada nos capítulos seguintes. Por enquanto, vamos dar por certa a arimetização da matemática, embora essa tenha sido uma proeza de importância capital.

Tendo reduzido toda a matemática pura tradicional à teoria dos números naturais, o passo seguinte em análise lógica foi reduzir essa teoria ela própria ao menor conjunto de premissas e termos indefinidos de que era possível derivá-la. Esse trabalho foi levado a cabo por Peano. Ele mostrou que toda a teoria dos números naturais podia ser derivada de três idéias primitivas e cinco proposições primitivas além daquelas da lógica pura. Essas três idéias e cinco proposições tornaram-se dessa maneira, por assim dizer, reféns de toda a matemática pura tradicional. Se elas pudessem ser definidas e provadas em termos de outras, toda a matemática pura também poderia sê-lo. Seu “peso” lógico, se podemos usar esse termo, é igual ao de toda a série de ciências que foram deduzidas da teoria dos números naturais; a verdade dessa série toda é assegurada se a verdade das cinco proposições primitivas estiver garantida, contanto, é claro, que não haja nada errôneo no aparato puramente lógico que também está aí envolvido. O trabalho de análise matemática é extraordinariamente facilitado por esse trabalho de Peano.

As três idéias primitivas na aritmética de Peano são:

0, número, sucessor.

Por “sucessor”, ele entende o número seguinte na ordem natural. Isto é, o sucessor de 0 é 1, o sucessor de 1 é 2, e assim por diante. Por “número” ele entende, nessa conexão, a classe dos números naturais.² Ele não está

supondo que conhecemos todos os membros dessa classe, apenas que sabemos o que temos em mente quando dizemos que isso ou aquilo é um número, assim como sabemos o que temos em mente quando dizemos “João é um homem”, embora não conheçamos todos os homens individualmente.

As cinco proposições primitivas que Peano supõe são:

- (1) 0 é um número.
- (2) O sucessor de qualquer número é um número.
- (3) Dois números diferentes nunca têm o mesmo sucessor.
- (4) 0 não é o sucessor de nenhum número.
- (5) Qualquer propriedade que pertença a 0 e também ao sucessor de qualquer número que tenha essa propriedade pertence a todos os números.

Este último é o princípio da indução matemática. Teremos muito a dizer com relação à indução matemática adiante; por enquanto, ela nos interessa unicamente tal como ocorre na análise da aritmética feita por Peano.

Consideremos brevemente de que maneira a teoria dos números naturais resulta dessas três idéias e cinco proposições. Para começar, definimos 1 como “o sucessor de 0”, 2 como “o sucessor de 1”, e assim por diante. Podemos obviamente ir tão longe quanto queiramos com essas definições, uma vez que, em virtude de (2), todo número a que chegemos terá um sucessor, e, em virtude de (3), esse não pode ser nenhum dos números já definidos, porque, se fosse, dois números diferentes teriam o mesmo sucessor; e em virtude de (4) nenhum dos números a que chegemos na série de sucessores pode ser 0. Assim, a série de sucessores nos dá uma série interminável de números continuamente novos. Em virtude de (5), todos os números pertencem a essa série, que começa com 0 e prossegue através de seus sucessivos sucessores: pois (a) 0 pertence a essa série, e (b) se um número n pertence a ela, seu sucessor também pertence; portanto, por indução matemática, todo número pertence a essa série.

Suponhamos que desejemos definir a soma de dois números. Tomando qualquer número m , definimos $m + 0$ como m , e $m + (n + 1)$ como o sucessor de $m + n$. Em virtude de (5), isso dá uma definição da soma de m e n , não importa que número n seja. Podemos definir de maneira semelhante o produto de dois números quaisquer. O leitor pode se convencer facilmente de que qualquer proposição elementar comum da

aritmética pode ser provada por meio de nossas cinco premissas, e se tiver alguma dificuldade poderá encontrar a prova em Peano.

Chegou o momento de nos voltarmos para considerações que tornam necessário avançar, além do ponto de vista de Peano — que representa a perfeição máxima da “arimetização” da matemática —, para o de Frege, que foi o primeiro a conseguir “logicizar” a matemática, isto é, reduzir à lógica as noções aritméticas que seus predecessores haviam demonstrado ser suficientes para a matemática. Não daremos realmente, neste capítulo, a definição de número e de números particulares de Frege, mas apresentaremos algumas das razões pelas quais o tratamento de Peano é menos definitivo do que parece.

Em primeiro lugar, as três idéias primitivas de Peano — a saber, “0”, “número” e “sucessor” — são passíveis de infinitas interpretações diferentes, todas as quais satisfarão as cinco proposições primitivas. Daremos alguns exemplos.

(1) Suponhamos que “0” significa 100 e que “número” seja tomado como significando os números de 100 em diante na série dos números naturais. Nesse caso, todas as nossas proposições primitivas ficam atendidas, mesmo a quarta, pois, embora 100 seja o sucessor de 99, 99, não é um “número” no sentido que estamos dando agora à palavra “número”. É óbvio que qualquer número pode substituir 100 neste exemplo.

(2) Suponhamos que “0” tem seu sentido usual, mas que “número” significa o que chamamos usualmente de “números pares”, e suponhamos que o “sucessor” de um número é o que resulta da adição de dois. Nesse caso, “1” será substituído pelo número dois, “2” será substituído pelo número quatro, e assim por diante; a série de “números” agora será

0, dois, quatro, seis, oito . . .

Todas as cinco premissas de Peano continuam a ser satisfeitas.

(3) Suponhamos que “0” significa o número um e que “número” significa o conjunto

$1, \frac{1}{2}, \frac{1}{4}, \frac{1}{8}, \frac{1}{16}, \dots$

e suponhamos que “sucessor” significa “metade”. Nesse caso todos os cinco axiomas de Peano se aplicarão a esse conjunto.

É claro que os exemplos poderiam ser multiplicados indefinidamente. De fato, dada qualquer série

$$x_0, x_1, x_2, x_3, \dots, x_n, \dots$$

que é interminável, não contém repetição alguma, tem um começo e não tem nenhum termo que não possa ser alcançado a partir do começo num número finito de passos, temos um conjunto de termos que verificam o axioma de Peano. Isso pode ser visto facilmente, embora a prova formal seja um pouco longa. Suponhamos que “0” significa x_0 , e que “número” significa todo o conjunto de termos, e suponhamos que o “sucessor” de x_n significa x_{n+1} . Então

- (1) “0 é um número”, isto é, x_0 é membro do conjunto.
 - (2) “O sucessor de qualquer número é um número”, isto é, tomando qualquer termo x_n no conjunto, x_{n+1} está também no conjunto
 - (3) “Dois números diferentes nunca têm o mesmo sucessor”, isto é, se x_m e x_n são dois membros diferentes do conjunto, x_{m+1} e x_{n+1} são diferentes; isso resulta do fato de que (por hipótese) não há repetições no conjunto.
 - (4) “0 não é o sucessor de nenhum número”, isto é, nenhum termo no conjunto vem antes de x_0 .
 - (5) Isso se torna: qualquer propriedade que pertença a x_0 e pertença a x_{n+1} , contanto que pertença a x_n , pertence a todos os x 's.
- Isso se segue da propriedade correspondente para números.
Uma série da forma

$$x_0, x_1, x_2, \dots, x_n, \dots$$

em que há um primeiro termo, um sucessor de cada termo (de modo que não há um último termo), nenhuma repetição, e cada termo pode ser alcançado a partir do começo num número finito de passos, é chamada *progressão*. Progressões são de grande importância nos princípios da matemática. Como acabamos de ver, toda progressão verifica os cinco axiomas de Peano. Inversamente, pode ser provado que toda série que verifica os cinco axiomas de Peano é uma progressão: “progressões” são “aquelas séries que verificam os cinco axiomas”. Qualquer progressão pode ser tomada como a base da matemática pura: podemos dar o nome

“0” ao primeiro termo, o nome “número” a todo o conjunto de seus termos, e o nome “sucessor” ao número seguinte na progressão. A progressão não precisa ser composta de números: pode ser composta de pontos no espaço, ou de momentos de tempo, ou de quaisquer outros termos de que haja uma provisão infinita. Cada diferente progressão dará origem a uma diferente interpretação de todas as proposições da matemática pura tradicional; todas essas possíveis interpretações serão igualmente verdadeiras.

No sistema de Peano, nada há que nos permita distinguir entre essas diferentes interpretações de suas idéias primitivas. Presume-se que sabemos o que se quer dizer por “0”, e que não devemos supor que esse símbolo significa 100, ou obelisco, ou qualquer das outras coisas que poderia significar.

Esse ponto, isto é, que “0” e “número” e “sucessor” não podem ser definidos por meio dos cinco axiomas de Peano, devendo ser compreendidos independentemente, é importante. Queremos que nossos números não meramente verifiquem fórmulas matemáticas, mas que se apliquem da maneira correta a objetos comuns. Queremos ter dez dedos, dois olhos e um nariz. Um sistema em que “1” significasse 100 e “2” significasse 101, e assim por diante, poderia estar muito bem para a matemática pura, mas não seria apropriado à vida diária. Queremos que “0” e “número” e “sucessor” tenham significados que nos dêem a quantidade certa de dedos, olhos e narizes. Já temos algum conhecimento (embora não suficientemente distinto ou analítico) do que queremos dizer por “1” e “2” e assim por diante, e nosso uso dos números na aritmética deve se conformar a esse conhecimento. Não podemos assegurar que esse será o caso usando-se o método de Peano; a única coisa que podemos dizer, se adotamos esse método, é que “sabemos o que queremos dizer por ‘0’ e ‘número’ e ‘sucessor, embora não sejamos capazes de explicar o que temos em mente em termos de outros conceitos mais simples”. É perfeitamente legítimo dizer isso quando *precisamos*, e em *algum* momento todos precisamos; mas o objetivo da filosofia matemática é não dizer isso por tanto tempo quanto possível. Pela teoria lógica da aritmética somos capazes de adiar isso por um tempo muito longo.

Seria possível sugerir que, em vez de estabelecer “0”, “número” e “sucessor” como termos cujo significado conhecemos, embora não sejamos capazes de defini-los, poderíamos deixá-los representar *quaisquer*

três termos que verifiquem os cinco axiomas de Peano. Nesse caso eles deixarão de ser termos que têm um significado definido, embora este não seja definido: serão termos “variáveis” com relação aos quais fazemos certas hipóteses, a saber, aquelas expressas nos cinco axiomas, mas que são sob outros aspectos indeterminados. Se adotarmos esse plano, nossos teoremas não serão provados com relação a um conjunto determinado de termos chamado “os números naturais”, mas com relação a todos os conjuntos de termos que possuam certas propriedades. Semelhante procedimento não é falacioso; de fato, para certos propósitos, representa uma generalização valiosa. De dois pontos de vista, porém, ele é incapaz de fornecer uma base adequada para a aritmética. Em primeiro lugar, não nos capacita a saber se há algum conjunto de termos que verifique os axiomas de Peano; não dá sequer a mais leve sugestão de alguma maneira de se descobrir se conjuntos desse tipo existem. Em segundo lugar, como já foi observado, queremos que nossos números sejam tais que possam ser usados para contar objetos comuns, e isso requer que nossos números possuam um significado *definido*, não que meramente possuam certas propriedades formais. Esse significado definido é definido pela teoria lógica da aritmética.

Capítulo 2

Definição de número

A pergunta “Que é um número?” foi formulada muitas vezes, mas só foi corretamente respondida em nosso próprio tempo. A resposta foi dada por Frege em 1884, em seu *Grundlagen der Arithmetik*.¹ Embora seja bastante pequeno, nada difícil e da mais elevada importância, esse livro não chamou quase nenhuma atenção, e a definição de número que contém permaneceu praticamente desconhecida até ser redescoberta pelo presente autor em 1901.

Ao procurar uma definição de número, a primeira coisa sobre a qual precisamos ter clareza é o que podemos chamar a de gramática de nossa investigação. Muitos filósofos, quando tentam definir número, estão na verdade empenhados em definir pluralidade, o que é uma coisa inteiramente diferente. *Número* é o que é característico dos números, como *homem* é o que é característico dos homens. Uma pluralidade não é um caso de número, mas de algum número particular. Um trio de homens, por exemplo, é um caso do número 3, e o número 3 é um caso de número; mas o trio não é um caso de número. Esse ponto pode parecer elementar e quase nem merecer menção; no entanto, provou-se demasiado sutil para os filósofos, com poucas exceções.

Um número particular não é idêntico a nenhuma coleção de termos que possua esse número: o número 3 não é idêntico ao trio que consiste em Brown, Jones e Robinson. O número 3 é algo que todos os trios têm em comum e que os distingue de outras coleções. Um número é algo que caracteriza certas coleções, a saber, aquelas que têm aquele número.

Em vez de falar de uma “coleção”, falaremos em geral de uma “classe”, ou por vezes de um “conjunto”. Outras palavras usadas em matemática para a mesma coisa são “agregado” e “múltiplo”. Teremos muito a dizer mais tarde sobre classes. Por ora, diremos tão pouco quanto possível. Mas há algumas observações que devem ser feitas imediatamente.

Uma classe ou coleção pode ser definida de duas maneiras que, à primeira vista, parecem muito diferentes. Podemos enumerar seus membros, como quando dizemos “a coleção a que me refiro é Brown, Jones e Robinson”. Ou podemos mencionar uma propriedade definidora, como quando falamos de “humanidade” ou “os habitantes de Londres. A definição que enumera é chamada uma definição por “extensão”, e a que menciona uma propriedade definidora é chamada uma definição por “intensão”. Desses dois tipos de definição, aquele por intensão é logicamente o mais fundamental. Isso é demonstrado por duas considerações: (1) que a definição extensional pode sempre ser reduzida a uma intensional; (2) que a definição intensional muitas vezes não pode nem teoricamente ser reduzida à extensional. Cada um desses pontos requer uma palavra de explicação.

(1) Brown, Jones e Robinson possuem todos eles certa propriedade que não é possuída por nada mais em todo o universo, a saber, a propriedade de ser ou Brown, ou Jones ou Robinson. Essa propriedade pode ser usada para dar uma definição por intensão da classe que consiste em Brown, Jones e Robinson. Considere uma fórmula como “ x é Brown ou x é Jones ou x é Robinson”. Essa fórmula será verdadeira para apenas três x , a saber, Brown, Jones e Robinson. Sob esse aspecto, ela se assemelha a uma equação cúbica com suas três raízes. Pode ser tomada como atribuindo uma propriedade comum aos membros da classe que consiste nesses três homens, e peculiar a eles. Um tratamento similar pode obviamente ser aplicado a qualquer outra classe dada em extensão.

(2) É óbvio que, na prática, podemos freqüentemente saber muito sobre uma classe sem sermos capazes de enumerar seus membros. De fato, nenhum homem poderia enumerar todos os homens, ou mesmo todos os habitantes de Londres; no entanto, sabe-se muito sobre cada uma dessas classes. Isso basta para mostrar que a definição por extensão não é necessária para o conhecimento acerca de uma classe. Mas quando passamos a considerar classes infinitas, descobrimos que a enumeração não é sequer teoricamente possível para seres que vivem apenas por um tempo finito. Não podemos enumerar todos os números naturais: eles são 0, 1, 2, 3, *e assim por diante*. Em algum ponto devemos nos contentar com “e assim por diante”. Não podemos enumerar todas as frações ou todos os números irracionais, ou todos os termos de qualquer outra coleção infinita.

Assim, nosso conhecimento acerca de todas essas coleções só pode ser derivado de uma definição por intensão.

Quando estamos procurando a definição de número, essas observações são relevantes de três diferentes maneiras. Em primeiro lugar, os próprios números formam uma coleção infinita, e não podem, portanto, ser definidos por enumeração. Em segundo lugar, as coleções que têm um número dado de termos formam elas próprias, presumivelmente, uma coleção infinita: deve-se presumir, por exemplo, que há uma coleção infinita de trios no mundo, pois se esse não fosse o caso o número total de coisas no mundo seria finito, o que, embora possível, parece improvável. Em terceiro lugar, desejamos definir “número” de tal maneira que números infinitos possam ser possíveis; assim devemos ser capazes de falar do número de termos em uma coleção infinita, e semelhante coleção deve ser definida por intensão, isto é, por uma propriedade comum a todos os seus membros e peculiar a eles.

Para muitos propósitos, uma classe e uma característica que a defina são praticamente intercambiáveis. A diferença vital entre as duas consiste no fato de que há somente uma classe que possui determinado conjunto de membros, ao passo que há sempre muitas diferentes características pelas quais uma dada classe pode ser definida. Os homens podem ser definidos como bípedes implumes, ou como animais racionais, ou (mais corretamente) pelos traços com que Swift delinea os Yahoos. É esse fato de uma característica definidora nunca ser única que torna as classes úteis; de outro modo poderíamos nos contentar com as propriedades comuns e peculiares a seus membros.² Qualquer uma dessas propriedades pode ser usada em lugar da classe sempre que a singularidade não for importante.

Retornando agora à definição de número, é claro que número é uma maneira de reunir certas coleções, a saber, aquelas que têm um dado número de termos. Podemos supor todos os pares num feixe, todos os trios em outro, e assim por diante. Dessa maneira, obtemos vários feixes de coleções, cada feixe consistindo em todas as coleções que têm certo número de termos. Cada feixe é uma classe cujos membros são coleções isto é, classes; assim, cada um é uma classe de classes. O feixe que consiste em todos os pares, por exemplo, é uma classe de classes: cada par é uma classe com dois membros, e todo o feixe de pares é uma classe com um número infinito de membros, cada um dos quais é uma classe de dois membros.

Como decidiremos se duas coleções devem pertencer ao mesmo feixe? A resposta que se sugere a si mesma é: “Descubra quantos membros cada uma tem, e ponha-as no mesmo feixe se tiverem o mesmo número de membros.” Mas isso pressupõe que temos números definidos, e que sabemos como descobrir quantos termos tem uma coleção. Estamos tão acostumados com a operação de contar que semelhante pressuposição poderia facilmente passar despercebida. De fato, no entanto, contar é uma operação, embora conhecida, logicamente muito complexa; além do mais ela só é disponível como um meio de descobrir quantos termos tem uma coleção quando a coleção é finita. Nossa definição de número não deve supor de antemão que todos os números são finitos; e não podemos de maneira alguma, sem um círculo vicioso, usar a contagem para definir números, porque usamos números para contar. Precisamos, portanto, de algum outro método para decidir quando duas coleções têm o mesmo número de termos.

Na realidade, é mais simples logicamente descobrir se duas coleções têm o mesmo número de termos do que definir que número é esse. Uma ilustração deixará isso claro. Se não houvesse nenhuma poligamia ou poliandria em lugar algum no mundo, é claro que o número de maridos vivos em determinado momento seria exatamente igual ao número de esposas. Não precisamos de um censo para nos assegurar disso, tampouco precisamos saber qual é o número real de maridos e esposas. Sabemos que o número deve ser o mesmo em ambas as coleções porque cada marido tem uma esposa e cada esposa tem um marido. A relação de marido e esposa é chamada “um-um”.

Diz-se que uma relação é “um-um” quando, se x tem a relação em questão com y , nenhum outro termo x' tem a mesma relação com y , e x não tem a mesma relação com nenhum termo y' que não y . Quando somente a primeira dessas duas condições é preenchida, a relação é chamada “um-muitos”; quando somente a segunda é preenchida, ela é chamada “muitos-um”. Convém observar que o número 1 não é usado nessas definições.

Nos países cristãos, a relação de marido para esposa é um-um; em países maometanos é um-muitos; no Tibéte é muitos-um.³ A relação de pai para filho é um-muitos; a relação de filho para pai é muitos-um, mas a de filho mais velho para pai é um-um. Se n for algum número, a relação de n para $n + 1$ é um-um; assim também é a relação de n para $2n$ ou $3n$. Quando

estamos considerando apenas números positivos, a relação de n para n^2 é um-um; mas quando números negativos são admitidos, ela se torna dois-um, pois n e $-n$ têm o mesmo quadrado. Esses exemplos deveriam ser suficientes para deixar claras as noções de relações um-um, um-muitos, muitos-um, que desempenham importante papel nos princípios da matemática, não somente em relação à definição de número, mas em muitas outras conexões.

Diz-se que duas classes são “similares” quando há uma relação um-um que correlaciona cada um dos termos de uma com um termo da outra, da mesma maneira como a relação de casamento correlaciona maridos com esposas. Algumas definições preliminares nos ajudarão a expressar essa definição de maneira mais precisa. A classe daqueles termos que têm uma dada relação com uma coisa ou outra é chamada o *domínio* daquela relação: assim pais são o domínio da relação de pai para filho, maridos são o domínio da relação de marido para esposa, esposas são o domínio da relação de esposa para marido, e maridos e esposas juntos são o domínio da relação de casamento. A relação de esposa para marido é chamada o *inverso* da relação de marido para esposa. De maneira similar, *menor* é o inverso de *maior*, *depois* é o inverso de *antes*, e assim por diante. Em geral, o inverso de uma dada relação é a relação que existe entre y e x sempre que a relação dada existe entre x e y . O *domínio inverso* de uma relação é o domínio de seu inverso: assim a classe das esposas é o domínio inverso da relação de marido para esposa. Podemos agora expressar nossa definição de similaridade da seguinte maneira:

Diz-se que uma classe é “similar” a outra quando há uma relação um-um da qual a primeira é o domínio, enquanto a outra é o domínio inverso.

É fácil provar (1) que toda classe é similar a si mesma; (2) que se uma classe α é similar a uma classe β , então β é similar a α ; (3) que se α é similar a β e β a γ , então α é similar a γ . Diz-se que uma relação é *reflexiva* quando possui a primeira dessas propriedades, *simétrica* quando possui a segunda, e *transitiva* quando possui a terceira. É óbvio que uma relação simétrica e transitiva deve ser reflexiva em todo o seu domínio. Relações que possuem essas propriedades são um tipo importante, e vale a pena observar que a similaridade é uma das relações desse tipo.

É óbvio ao senso comum que duas classes finitas têm o mesmo número de termos se forem similares, mas não de outro modo. O ato de contar consiste em estabelecer uma correlação um-um entre o conjunto de

objetos contados e os números naturais (excluindo 0) que são usados no processo. Dessa maneira, o senso comum conclui que há tantos objetos no conjunto a ser contado quantos são os números até o último número a ser usado na contagem. E sabemos também que, enquanto nos limitarmos aos números finitos, há exatamente n números de 1 até n . Segue-se, portanto, que o último número usado na contagem de uma coleção é o número de termos na coleção, contanto que esta seja finita. Mas esse resultado, além de ser aplicável apenas a coleções finitas, depende do fato de que duas classes que são similares têm o mesmo número de termos e pressupõe esse fato: pois o que fazemos quando contamos (digamos) dez objetos é mostrar que o conjunto desses objetos é similar ao conjunto dos números 1 a 10. A noção de similaridade é logicamente pressuposta na operação de contagem, e é logicamente mais simples, embora menos familiar. Ao contar, é necessário tomar os objetos contados numa certa ordem, como primeiro, segundo, terceiro etc., mas a ordem não é da essência do número: é uma adição irrelevante, uma complicação desnecessária do ponto de vista lógico. A noção de similaridade não exige uma ordem: por exemplo, vimos que o número de maridos é igual ao de esposas sem precisar estabelecer uma ordem de precedência entre eles. A noção de similaridade também não requer que as classes similares sejam finitas. Tomemos, por exemplo, os números naturais (excluindo 0) por um lado, e as frações cujo numerador é 1 por outro lado: é óbvio que podemos correlacionar 2 com $1/2$, 3 com $1/3$ e assim por diante, provando desse modo que as duas classes são similares.

Podemos, portanto, usar a noção de “similaridade” para decidir quando duas coleções devem pertencer ao mesmo feixe, no sentido em que formulamos essa questão anteriormente neste capítulo. Queremos fazer um feixe contendo a classe que não tem membros: este será para o número 0. Depois queremos um feixe para todas as classes que têm um membro: este será para o número 1. Depois, para o número 2, queremos um feixe que consista em todos os pares; depois um de todos os trios; e assim por diante. Dada qualquer coleção, podemos definir o feixe a que deve pertencer como sendo a classe de todas aquelas coleções que são “similares” a ela. É muito fácil ver que se (por exemplo) uma coleção tem três membros, a classe de todas as coleções que são similares a ela será a classe dos trios. E seja qual for o número de termos que uma coleção possa ter, aquelas coleções que são “similares” a ela terão o mesmo número de

termos. Podemos tomar isso como uma *definição* de “ter o mesmo número de termos”. É óbvio que ela fornece resultados compatíveis com o uso desde que nos limitemos a coleções finitas.

Até agora não sugerimos nada que seja minimamente paradoxal. Quando chegamos à definição real de números, porém, não podemos evitar o que deve à primeira vista parecer um paradoxo, embora essa impressão vá logo desaparecer. Naturalmente pensamos que a classe dos pares (por exemplo) é algo diferente do número 2. Mas não há dúvida com relação à classe dos pares: ela é indubitável e não difícil de definir, ao passo que o número 2, em qualquer outro sentido, é uma entidade metafísica com relação à qual nunca teremos certeza de que existe ou de que a apreendemos. É mais prudente, portanto, nos contentarmos com a classe dos pares, de que temos certeza, do que sairmos à caça de um problemático número 2 que deverá permanecer sempre elusivo.

Conseqüentemente, formulamos a seguinte definição:

O número de uma classe é a classe de todas as classes que são similares a ele.

O número de um par, portanto, será a classe de todos os pares. De fato, a classe de todos os pares será o número 2, segundo nossa definição. À custa de uma pequena esquisitice, essa definição assegura exatidão e indubitabilidade; e não é difícil provar que números assim definidos têm todas as propriedades que esperamos que os números tenham.

Podemos agora ir adiante para definir números em geral como qualquer dos feixes em que a similaridade reúna classes. Um número será um conjunto de classes tal que quaisquer duas são similares uma à outra, e nenhuma fora do conjunto é similar a alguma dentro dele. Em outras palavras, um número (em geral) é qualquer coleção que é o número de um de seus membros; ou, ainda mais simplesmente:

Um número é qualquer coisa que é o número de alguma classe.

Semelhante definição tem uma aparência verbal de ser circular, mas de fato não é. Definimos “o número de uma dada classe” sem usar a noção de número em geral; assim, podemos definir número em geral em termos de “o número de uma dada classe” sem cometer nenhum erro lógico.

Definições desse tipo são de fato muito comuns. A classe dos pais, por exemplo, teria de ser definida definindo-se primeiro o que é ser pai de alguém; depois a classe dos pais será todos os que são pais de alguém. De maneira similar, se queremos definir números quadrados (digamos), temos

primeiro de definir o que temos em mente ao dizer que um número é o quadrado de outro, e depois definir números quadrados como aqueles que são os quadrados de outros números. Esse tipo de procedimento é muito comum, e é importante compreender que é legítimo e com muita frequência necessário.

Demos agora uma definição de números que servirá para coleções finitas. Resta ver como servirá para coleções infinitas. Antes, porém, precisamos decidir o que entendemos por “finito” e “infinito”, o que não pode ser feito nos limites do presente capítulo.

Capítulo 3

Finitude e indução matemática

A série dos números naturais, como vimos no Capítulo 1, pode ser definida se soubermos o que queremos dizer pelos três termos “0”, “número” e “sucessor”. Mas vamos dar um passo adiante: podemos definir todos os números naturais se soubermos o que queremos dizer por “0” e “sucessor”. Será útil compreender a diferença entre finito e infinito para ver como isso pode ser feito e por que o método pelo qual é feito não pode ser estendido além do finito. Ainda não vamos considerar de que maneira “0” e “sucessor” devem ser definidos: por enquanto vamos supor que sabemos o que esses termos significam e mostrar como, a partir deles, todos os outros números naturais podem ser obtidos.

É fácil ver que podemos alcançar qualquer número designado, digamos 30.000. Em primeiro lugar definimos “1” como “o sucessor de 0”, depois definimos 2 como “o sucessor de 1”, e assim por diante. No caso de um número designado, como 30.000, a prova de que podemos alcançá-lo procedendo passo a passo dessa maneira pode ser obtida, se tivermos paciência, por experimentação real: podemos avançar até realmente chegarmos a 30.000. Embora o método da experimentação esteja disponível para cada número natural particular, não está disponível para provar a proposição geral de que *todos* esses números podem ser alcançados dessa maneira, isto é, procedendo-se passo a passo de cada número para seu sucessor. Haverá alguma outro modo pela qual isso possa ser provado?

Consideremos a questão pelo lado oposto. Quais são os números que podem ser alcançados, dados os termos “0” e “sucessor”? Há alguma maneira pela qual possamos definir toda a classe desses números? Alcançamos 1, como o sucessor de 0; 2, como o sucessor de 1; 3, como o sucessor de 2; e assim por diante. É esse “e assim por diante” que desejamos substituir por algo menos vago e indefinido. Poderíamos ser

tentados a dizer que “e assim por diante” significa que o processo de avançar para o sucessor pode ser repetido *qualquer número finito* de vezes; mas o problema em que estamos envolvidos é o de definir “número finito”, e portanto não devemos usar essa noção em nossa definição. Ela não deve presumir que sabemos o que é número finito.

A chave para nosso problema reside na *indução matemática*. Lembremos que, no Capítulo 1, essa era a quinta das cinco proposições primitivas que formulamos sobre os números naturais. Ela declarava que qualquer propriedade que pertença a 0 e ao sucessor de qualquer número que a possua pertence a todos os números naturais. Isso foi apresentado então como um princípio, mas devemos agora adotá-lo como uma definição. Não é difícil ver que os termos que obedecem a ela são os mesmos que podem ser alcançados a partir de 0 por passos sucessivos de um número para o seguinte, mas como o ponto é importante, exporemos a matéria com alguns detalhes.

Convém começar com algumas definições, que serão úteis também em outras conexões.

Diz-se que uma propriedade é “hereditária” na série dos números naturais se, sempre que ela pertence a um número n , pertence também a $n + 1$, o sucessor de n . De maneira similar, diz-se que uma classe é “hereditária” se, sempre que n é membro da classe, $n + 1$ também é. É fácil ver, embora ainda não se possa presumir que saibamos isso, que dizer que uma propriedade é hereditária equivale a dizer que ela pertence a todos os números naturais não menores do que algum deles, por exemplo, ela deve pertencer a todos não menores que do 100, ou a todos não menores do que 1.000, ou pode ser que pertença a todos não menores do que 0, isto é, a todos sem exceção.

Diz-se que uma propriedade é “indutiva” quando é uma propriedade hereditária que pertence a 0. De maneira similar, uma classe é indutiva quando é uma classe hereditária de que 0 é membro.

Dada uma classe hereditária de que 0 é membro, segue-se que 1 é membro dela, porque uma classe hereditária contém o sucessor de seus membros, e 1 é o sucessor de 0. De maneira similar, dada uma classe hereditária de que 1 é membro, segue-se que 2 é membro dela; e assim por diante. Podemos provar desse modo, por um procedimento passo a passo, que qualquer número natural designado, digamos 30.000, é membro de toda classe indutiva.

Definiremos a “posteridade” de um dado número natural com respeito à relação “predecessor imediato” (que é o inverso de “sucessor”) como todos aqueles termos que pertencem a toda classe hereditária a que o número dado pertence. Novamente, é fácil *ver* que a posteridade de um número natural consiste nele mesmo e todos os números naturais maiores; mas também isto nós ainda não sabemos oficialmente.

Pelas definições mencionadas, a posteridade de 0 consistirá naqueles termos que pertencem a toda classe indutiva.

Não é difícil agora deixar óbvio que a posteridade de 0 é aquele mesmo conjunto composto por aqueles termos que podem ser alcançados a partir de 0 por passos sucessivos de um para o seguinte. Pois, em primeiro lugar, 0 pertence a ambos esses conjuntos (no sentido em que definimos nossos termos); em segundo lugar, se n pertence a ambos os conjuntos, $n + 1$ também pertence. É preciso observar que estamos tratando aqui do tipo de matéria que não admite prova precisa, a saber, a comparação de uma idéia relativamente vaga com outra relativamente precisa. A noção de “aqueles termos que podem ser alcançados a partir de 0 por passos sucessivos de um para o seguinte” é vaga, embora *pareça* transmitir um significado definido; por outro lado, “a posteridade de 0” é uma noção precisa e explícita, exatamente onde a outra idéia é nebulosa. Ela pode ser tomada como expressando o que *tínhamos* em mente quando falamos dos termos que podem ser alcançados a partir de 0 por passos sucessivos.

Formulamos agora a seguinte definição:

Os “números naturais” são a posteridade de 0 com respeito à relação “predecessor imediato” (que é o inverso de “sucessor”).

Chegamos assim a uma definição de uma das três idéias primitivas de Peano em termos das outras duas. Em resultado dessa definição, duas das proposições primitivas dele — a saber, a que afirma que 0 é um número e a que afirma a indução matemática — se tornam desnecessárias, pois resultam da definição. Aquela que afirma que o sucessor de um número natural é um número natural é necessária apenas na forma enfraquecida “todo número natural tem um sucessor”.

Podemos, é claro, definir facilmente “0” e “sucessor” por meio da definição de número em geral a que chegamos no Capítulo 2. O número 0 é o número de termos numa classe que não tem membro algum, isto é, na classe chamada a “classe nula”. Pela definição geral de número, o número de termos na classe nula é o conjunto de todas as classes similares à classe

nula, isto é (como pode ser facilmente provado), o conjunto que consiste unicamente na classe nula, ou seja, a classe cujo único membro é a classe nula. (Esta não é idêntica à classe nula: ela tem um membro, a saber, a classe nula, ao passo que a própria classe nula não tem membro algum, como explicaremos quando chegarmos à teoria das classes.) Temos, portanto, a seguinte definição puramente lógica:

0 é a classe cujo único membro é a classe nula.

Resta definir “sucessor”. Dado qualquer número n , suponhamos que α é uma classe que tem n membros, e que x é um termo que não é membro de α . Portanto, a classe que consiste em α com o acréscimo de x terá $n + 1$ membros. Temos assim a seguinte definição:

O sucessor do número de termos na classe α é o número de termos na classe que consiste em α juntamente com x , em que x é qualquer termo não pertencente à classe.

Certos refinamentos são necessários para tornar essa definição perfeita, mas não precisamos nos preocupar com eles.¹ Deve-se lembrar que já demos (no capítulo 2) uma definição lógica do número de termos numa classe, a saber, o definimos como o conjunto de todas as classes similares à classe dada.

Reduzimos assim as três idéias primitivas de Peano a idéias de lógica: demos definições delas que as tornam definidas, não mais passíveis de ter uma infinidade de diferentes sentidos, como eram quando determinadas apenas na medida em que obedeciam aos cinco axiomas de Peano. Nós as removemos do aparato de termos que devem ser meramente apreendidos, e aumentamos assim a articulação dedutiva da matemática.

No tocante às cinco proposições primitivas, já conseguimos tornar duas delas demonstráveis mediante nossa definição de “número natural”. Como ficam as três restantes? É muito fácil provar que 0 não é o sucessor de nenhum número, e que o sucessor de algum número é um número. Mas há uma dificuldade acerca da proposição primitiva que resta, a saber, “dois números diferentes nunca têm o mesmo sucessor”. A dificuldade não surge a menos que o número total dos indivíduos no universo seja finito; para dois números dados m e n , nenhum dos quais é o número total de indivíduos no universo, é fácil provar que não podemos ter $m + 1 = n + 1$, a menos que tenhamos $m = n$. Suponhamos, porém, que o número total de indivíduos no universo fosse (digamos) 10; nesse caso, não haveria nenhuma classe de 11 indivíduos e o número 11 seria a classe nula. Assim

também seria o número 12. Portanto, teríamos $11 = 12$; portanto o sucessor de 10 seria igual ao sucessor de 11, embora 10 não fosse igual a 11. Teríamos, portanto, dois números diferentes com o mesmo sucessor. Esse fracasso do terceiro axioma não pode ocorrer, contudo, se o número de indivíduos no mundo não for finito. Retornaremos a esse tópico mais adiante.²

Supondo que o número de indivíduos no universo não é finito, já conseguimos não só definir as três idéias primitivas de Peano, mas ver como provar suas cinco proposições por meio de idéias e proposições primitivas pertencentes à lógica. Disso decorre que toda a matemática pura, na medida em que é dedutível da teoria dos números naturais, é apenas um prolongamento da lógica. A extensão desse resultado àqueles ramos modernos da matemática não dedutíveis da teoria dos números naturais não oferece nenhuma dificuldade de princípio, como mostramos em outro lugar.³

O processo de indução matemática por meio do qual definimos os números naturais é passível de generalização. Definimos os números naturais como a “posteridade” de 0 com respeito à relação de um número com seu sucessor imediato. Se chamarmos essa relação N , qualquer número m terá essa relação com $m + 1$. Uma propriedade é “hereditária com respeito a N ”, ou simplesmente “ N -hereditária”, se, sempre que a propriedade pertencer a um número m , pertencer também a $m + 1$, isto é, ao número com que m tem a relação N . E diremos que um número n pertence à “posteridade” de m com respeito à relação N se n tiver todas as propriedades hereditárias- N que pertencem a m . Essas definições podem todas ser aplicadas a qualquer outra relação assim como a N . Dessa maneira, se R for uma relação qualquer, podemos formular as seguintes definições:⁴

Uma propriedade é chamada “ R -hereditária” quando, se pertencer a um termo x , e x tiver a relação R com y , pertencer também a y .

Uma classe é hereditária- R quando sua propriedade definidora for hereditária- R .

Diz-se que um termo x é “ R -ancestral” do termo y se y tiver todas as propriedades hereditárias- R que x tiver, contanto que x seja um termo que tenha a relação R com alguma coisa, ou com o qual alguma coisa tenha a relação R . (Isso é apenas para excluir casos triviais.)

A “ R -posteridade” de x são todos os termos de que x é um R -ancestral.

Construímos as definições previamente mencionadas de tal modo que se um termo for o ancestral de alguma coisa, ele é seu próprio ancestral e pertence à sua própria posteridade. Isso é apenas por conveniência.

Deve-se observar que se tomarmos para R a relação “pai”, “ancestral” e “posteridade”, teremos os sentidos usuais, a não ser pelo fato de que uma pessoa será incluída entre seus próprios ancestrais e posteridade. Fica imediatamente óbvio que “ancestral” deve ser passível de definição em termos de “pai”, mas até que Frege desenvolvesse sua teoria generalizada da indução, ninguém poderia ter definido “ancestral” precisamente em termos de “pai”. Uma breve consideração desse ponto servirá para mostrar a importância da teoria. Alguém que se confrontasse pela primeira vez com o problema de definir “ancestral” em termos de “pai” diria naturalmente que A é um ancestral de Z se, entre A e Z , houver certo número de pessoas, $B, C, \dots$, entre as quais B é filho de A , cada um é pai do seguinte, até o último, que é pai de Z . Mas essa definição não é adequada, a menos que acrescentemos que o número de termos intermediários deve ser finito. Tomemos, por exemplo, uma série como a seguinte:

$$-1, -\frac{1}{2}, -\frac{1}{4}, -\frac{1}{8}, \dots \frac{1}{8}, \frac{1}{4}, \frac{1}{2}, 1.$$

Temos aqui primeiro uma série de frações negativas sem fim e depois uma série de frações positivas sem começo. Devemos dizer que, nessa série, -1 é ancestral de 1 ? Será assim de acordo com a definição para iniciantes sugerida anteriormente, mas não será de acordo com qualquer definição que dê o tipo de idéia que desejamos definir. Para esse propósito, é essencial que o número de intermediários seja finito. Mas, como vimos, “finito” deve ser definido por meio de indução matemática, e é mais simples definir a relação ancestral de maneira geral imediatamente que defini-la apenas para o caso da relação de n com $n + 1$, e depois estendê-la para os outros casos. Aqui, como constantemente em outros lugares, começar pela generalidade, embora possa exigir mais reflexão de início, provará no final das contas economizar reflexão e aumentar a capacidade lógica.

No passado, o uso de indução matemática em demonstrações foi uma espécie de mistério. Parecia não se poder duvidar sensatamente de que ela era um método de prova válido, mas ninguém sabia ao certo por quê.

Alguns acreditavam que ela era realmente um caso de indução, no sentido em que a palavra é usada em lógica. Poincaré considerava-a um princípio da máxima importância, por meio do qual um número infinito de silogismos podia ser condensado num único argumento. Sabemos agora que todas essas idéias são errôneas, e que a indução matemática é uma definição, não um princípio. Há alguns números aos quais pode ser aplicada, e há outros (como veremos no Capítulo 8) a que não pode. Nós *definimos* os “números naturais” como aqueles aos quais provas por indução matemática podem ser aplicadas, isto é como aqueles que possuem todas as propriedades indutivas. Segue-se que tais provas podem ser aplicadas aos números naturais não em virtude de alguma intuição, axioma ou princípio misteriosos, mas como uma proposição puramente verbal. Se “quadrúpedes” são definidos como animais de quatro patas, disso se seguirá que animais de quatro patas são quadrúpedes; e o caso dos números que obedecem à indução matemática é exatamente similar.

Usaremos a expressão “números indutivos” para designar o mesmo conjunto de que falamos até agora como “números naturais”. A expressão “números indutivos” é preferível por servir de lembrete de que a definição desse conjunto de números é obtida a partir da indução matemática.

A indução matemática proporciona, mais do que qualquer outra coisa, a característica essencial pela qual se pode distinguir o finito do infinito. O princípio da indução matemática poderia ser expresso popularmente sob alguma forma do tipo “o que pode ser inferido de vizinho para vizinho pode ser inferido do primeiro para o último”.

Isso é verdade quando o número de passos intermediários entre o primeiro e o último é finito, não de outra maneira. Quem já observou um trem de carga começando a se mover, terá notado como o impulso é comunicado com um solavanco de um vagão para o seguinte, até que finalmente o último vagão esteja em movimento. Quando o trem é muito longo, o último vagão leva bastante tempo para se mover. Se o trem fosse infinitamente longo, haveria uma sucessão infinita de solavancos, e nunca chegaria o momento em que o trem inteiro estaria em movimento. Apesar disso, se houvesse uma série de vagões não maior do que a série de números indutivos (que, como veremos, é um caso dos menores dos infinitos), cada vagão começaria a se mover mais cedo ou mais tarde se a locomotiva perseverasse, ainda que fosse haver sempre outros vagões mais atrás que ainda não teriam começado a se mover. Essa imagem ajudará a

elucidar o argumento de vizinho para vizinho, e sua conexão com a finitude. Quando passamos aos números infinitos, em que os argumentos obtidos por indução matemática deixarão de ser válidos, as propriedades desses números ajudarão a elucidar, por contraste, o uso quase inconsciente que é feito da indução matemática no tocante a números finitos.

Capítulo 4

A definição de ordem

Já levamos nossa análise da série dos números naturais até o ponto em que obtivemos definições lógicas dos membros dessa série, de toda a classe de seus membros e da relação de um número com seu sucessor imediato. Devemos agora considerar o caráter *serial* dos números naturais na ordem $0, 1, 2, 3, \dots$. Em geral, pensamos nos números como estando nessa *ordem*, e é uma parte essencial do trabalho de analisar nossos dados procurar uma definição de “ordem” ou “série” em termos lógicos.

A noção de ordem tem enorme importância em matemática. Não só os números inteiros, mas também as frações racionais e todos os números reais têm uma ordem de magnitude, e essa é essencial para a maioria de suas propriedades matemáticas. A ordem dos pontos numa linha é essencial para a geometria; assim também é a ordem ligeiramente mais complicada das linhas através de um ponto num plano, ou de planos ao longo de uma linha. As dimensões, em geometria, são um desenvolvimento da ordem. A concepção de *limite*, que é subjacente a toda a matemática superior, é uma concepção serial. Há partes da matemática que não dependem da noção de ordem, mas são muito poucas em comparação com aquelas em que essa noção está envolvida.

Ao procurar uma definição de ordem, a primeira coisa a compreender é que nenhum conjunto de termos tem apenas *uma* ordem, a ponto de excluir outras. Um conjunto de termos tem todas as ordens de que é capaz. Algumas vezes uma ordem é tão mais conhecida e natural para nossos pensamentos que nos inclinamos a considerá-la *a* ordem daquele conjunto de termos; mas isso é um erro. Os números naturais — ou os números “indutivos”, como também os chamaremos — ocorrem a nós mais prontamente em ordem de magnitude, mas eles são passíveis de um número infinito de outros arranjos. Poderíamos, por exemplo, considerar primeiramente os números ímpares e depois os números pares; ou

primeiro 1, depois todos os números pares, depois todos os múltiplos ímpares de 3, depois todos os múltiplos de 5 mas não de 2 ou 3, depois todos os múltiplos de 7 mas não de 2 ou 3 ou 5, e assim por diante através de toda a série dos primos. Quando dizemos que podemos “arranjar” os números nessas várias ordens, essa é uma expressão imprecisa: o que realmente fazemos é voltar nossa atenção para certas relações entre os números naturais, que geram eles próprios tal e tal arranjo. Somos tão incapazes de “arranjar” os números naturais quanto de arranjar o céu estrelado; mas, assim como podemos perceber entre as estrelas fixas seja sua ordem de brilho ou sua distribuição no céu, assim também há várias relações entre os números que podem ser observadas e que dão origem a várias ordens diferentes entre eles, todas igualmente legítimas. E o que é verdade acerca de números é igualmente verdade acerca de pontos numa linha ou dos momentos do tempo: uma ordem é mais conhecida, mas outras são igualmente válidas. Poderíamos, por exemplo, tomar primeiro, numa linha, todos os pontos que têm coordenadas integrais, depois todos os que têm coordenadas racionais não-integrais, depois todos os que têm coordenadas racionais algébricas, e assim por diante, através de qualquer conjunto de complicações que desejamos. A ordem resultante será uma que os pontos da linha certamente têm, quer optemos por notá-la ou não; a única coisa arbitrária nas várias ordens de um conjunto de termos é nossa atenção, pois os próprios termos têm sempre todas as ordens de que são capazes.

Um resultado importante dessa consideração é que não devemos procurar a definição de ordem na natureza do conjunto de termos a serem ordenados, uma vez que um conjunto de termos tem muitas ordens. A ordem reside, não na *classe* de termos, mas na relação entre os membros da classe, com respeito à qual alguns aparecem como anteriores, e alguns, como posteriores. O fato de uma classe poder ter muitas ordens deve-se ao fato de poder haver muitas relações entre os membros de uma única classe. Que propriedades deve uma relação possuir para dar origem a uma ordem?

As características essenciais de uma relação que dará origem a uma ordem podem ser descobertas considerando-se que, com respeito a tal relação, devemos ser capazes de dizer, acerca de dois termos quaisquer na classe a ordenar, que um “precede” e o outro “segue”. Ora, para podermos

usar essas palavras tal como as compreenderíamos naturalmente, é preciso que a relação ordenadora tenha três propriedades:

(1) Se x preceder y , y não deverá também preceder x . Esta é uma característica óbvia do tipo de relação que conduz a séries. Se x for menor do que y , y não será também menor do que x . Se x for anterior no tempo a y , y não será também anterior a x . Por outro lado, relações que não dão origem a séries com frequência não têm essa propriedade. Se x for irmão ou irmã de y , y será irmão ou irmã de x . Se x for da mesma altura que y , y será da mesma altura que x . Se a altura de x for diferente da de y , a altura de y será diferente da de x . Em todos esses casos, quando a relação existe entre x e y , existe também entre y e x . Com relações seriais, porém, tal coisa não acontece. Uma relação que tenha essa primeira propriedade é chamada *assimétrica*.

(2) Se x preceder y e y preceder z , x deve preceder z . Isso pode ser ilustrado pelos mesmos exemplos de antes: *menor*, *anterior*, *à esquerda de*. Mas como exemplos de relações que *não* têm essa propriedade, somente dois de nossos três exemplos anteriores servirão. Se x for irmão ou irmã de y , e y de z , x pode não ser irmão ou irmã de z , pois x e z podem ser a mesma pessoa. O mesmo se aplica à diferença de altura, mas não à igualdade de altura, que tem nossa segunda propriedade, mas não a primeira. A relação “pai”, por outro lado, tem nossa primeira propriedade, mas não a segunda. Uma relação que tenha a segunda propriedade é chamada *transitiva*.

(3) Dados quaisquer dois termos da classe que deve ser ordenada, deve haver um que preceda e outro que siga. Por exemplo, de quaisquer dois números inteiros, ou frações, ou números reais, um é menor e o outro maior; mas isso não é verdade com relação a quaisquer dois números complexos. De quaisquer dois momentos no tempo, um deve ser anterior ao outro, mas isso não pode ser dito sobre eventos, que podem ser simultâneos. De dois pontos numa linha, um deve estar à esquerda do outro. Uma relação que possua essa terceira propriedade é chamada *conexa*.

Quando uma relação possui essas três propriedades, ela é do tipo que dá origem a uma ordem em meio aos termos entre os quais existe; e onde quer que haja uma ordem, é possível encontrar alguma relação que possua essas três propriedades, gerando-a.

Antes de ilustrar essa tese, introduziremos algumas definições.

(1) Diz-se que uma relação é uma alio-relativa,¹ ou que *está contida em diversidade* ou que *implica diversidade*, se nenhum termo tiver essa relação consigo mesmo. Assim, por exemplo, “maior”, “diferente em tamanho”, “irmão”, “marido”, “pai” são alio-relativas; mas “iguais”, “nascido dos mesmos pais”, “caro amigo” não são.

(2) O *quadrado* de uma relação é aquela relação que existe entre dois termos x e z quando há um termo intermediário y tal que a relação dada exista entre x e y e entre y e z . Assim “avô paterno” é o quadrado de “pai”, “maior por 2” é o quadrado de “maior por 1”, e assim por diante.

(3) O *domínio* de uma relação consiste que todos os termos em têm a relação com uma coisa ou outra, e o *domínio inverso* consiste em todos os termos com que uma coisa ou outra têm a relação. Essas palavras já foram definidas, mas são lembradas aqui no interesse da definição seguinte:

(4) O *campo* de uma relação consiste em seu domínio e domínio inverso juntos.

(5) Diz-se que uma relação *contém* outra ou *é implicada por* outra se existir sempre que a outra existir.

Veremos que uma relação *assimétrica* é o mesmo que uma relação cujo quadrado é uma anti-reflexiva. Ocorre freqüentemente que uma relação seja uma anti-reflexiva sem ser assimétrica, embora uma relação assimétrica seja sempre uma anti-reflexiva. Por exemplo, “cônjuge” é uma anti-reflexiva, mas é assimétrica, pois se x é cônjuge de y , y é cônjuge de x . Mas entre relações *transitivas*, todas as anti-reflexiva são assimétricas, bem como vice-versa;

A partir dessas definições, podemos ver que uma relação *transitiva* é uma relação implicada por seu quadrado, ou, como também dizemos, que “contém” seu quadrado. Assim, “ancestral” é transitivo, porque o ancestral de um ancestral é um ancestral; mas “pai” não é transitivo, porque o pai de um pai não é um pai. Uma transitiva anti-reflexiva é uma relação que contém seu quadrado e é contida em diversidade; ou, o que vem a dar no mesmo, uma relação cujo quadrado implica tanto essa relação quanto diversidade — porque quando uma relação é transitiva, ser assimétrica é equivalente a ser uma anti-reflexiva.

Uma relação é *conexa* quando, dados quaisquer dois termos diferentes de seu campo, a relação existe entre o primeiro e o segundo ou entre o segundo e o primeiro (sem excluir a possibilidade de as duas coisas

ocorrerem, embora ambas não possam ocorrer se a relação for assimétrica.)

Veremos que a relação “ancestral”, por exemplo, é uma anti-reflexiva e transitiva, mas não conexa; é por não ser conexa que não é suficiente para arranjar a raça humana numa série.

A relação “menor do que ou igual a”, entre números, é transitiva e conexa, mas não assimétrica nem uma anti-reflexiva.

A relação “maior ou menor” entre números é uma anti-reflexiva e é conexa, mas não é transitiva, pois se x for maior ou menor do que y , e y for maior ou menor do que z , pode acontecer que x e z sejam o mesmo número.

Assim as três propriedades de ser (1) uma anti-reflexiva, (2) transitiva e (3) conexa são mutuamente independentes, visto que uma relação pode possuir quaisquer duas delas sem ter a terceira.

Formulamos agora a seguinte definição:

Uma relação é *serial* quando é uma anti-reflexiva, transitiva e conexa; ou, o que é equivalente, quando é assimétrica, transitiva e conexa.

Uma *série* é o mesmo que uma relação serial.

Seria possível pensar que uma série deveria ser o *campo* de uma relação serial, não a própria relação. Mas isso seria um erro. Por exemplo,

1, 2, 3; 1, 3, 2; 2, 3, 1; 2, 1, 3; 3, 1, 2; 3, 2, 1

são seis diferentes series que têm todas o mesmo campo. Se o campo *fosse* a série, só poderia haver uma única série com um dado campo. O que distingue as seis séries citadas são simplesmente as diferentes relações de ordenação nos seis casos. Dada a relação de ordenação, o campo e a ordem são ambos determinados. Assim a relação de ordenação pode ser tomada como *sendo* a série, mas o campo não pode ser assim tomado.

Dada qualquer relação serial, digamos P , diremos que, com respeito a essa relação, x “precede” y se x tiver a relação P com y , que escreveremos “ xPy ” para abreviar. As três características que P deve possuir para ser serial são:

(1) Não devemos nunca ter xPx , isto é, nenhum termo deve preceder a si mesmo.

(2) P^2 deve implicar P , isto é, se x precede y e y precede z , z deve preceder P .

(3) Se x e y forem dois termos diferentes no campo de P , teremos xPy ou yPx , isto é, um dos dois deve preceder o outro.

O leitor pode se convencer facilmente de que, quando essas três propriedades são encontradas numa relação de ordenação, as características que esperamos das séries poderão também ser encontradas, e vice-versa. É legítimo, portanto, tomar o que foi mencionado como uma definição de ordem ou série. Cabe ressaltar que a definição foi levada a cabo em termos puramente lógicos.

Embora sempre exista uma relação conexa assimétrica transitiva onde quer que haja uma série, ela não é sempre a relação que seria mais naturalmente vista como gerando a série. A série dos números naturais pode servir de ilustração. A relação que demos por certa ao considerar os números naturais foi a relação de sucessão imediata, isto é, a relação entre números inteiros consecutivos. Essa relação é assimétrica, mas não transitiva ou conexa. Podemos, contudo, derivar dela, pelo método da indução matemática, a relação “ancestral” que consideramos no capítulo anterior. Essa relação será o mesmo que “menor do que ou igual a” entre números inteiros indutivos. Para o objetivo de gerar a série de números naturais, queremos a relação “menor do que”, excluindo “igual a”. Essa é a relação de m com n quando m é um ancestral de n , mas não idêntico a n , ou (o que vem a dar no mesmo) quando o sucessor de m é um ancestral de n no sentido em que um número é seu próprio ancestral. Isso significa que formularemos a seguinte definição:

Dizemos que um número indutivo m é menor do que um outro número n quando n possui todas as propriedades hereditárias possuídas pelo sucessor de m .

É fácil ver, e não é difícil provar, que a relação “menor do que”, assim definida, é assimétrica, transitiva e conexa, e tem os números indutivos para seu campo. Assim, por meio dessa relação, os números indutivos adquirem uma ordem no sentido em que definimos o termo “ordem” e essa ordem é a chamada “ordem natural”, ou ordem de magnitude.

A geração de séries por meio de relações mais ou menos semelhantes àquela de n com $n + 1$ é muito comum. A série dos reis da Inglaterra, por exemplo, é gerada por relações de cada um com seu sucessor. Esse é

provavelmente o meio mais fácil, onde é aplicável, de conceber a geração de uma série. Nesse método, passamos de cada termo para o seguinte, enquanto houver um seguinte, ou de volta para o anterior, enquanto houver um anterior. Esse método sempre requer a forma generalizada da indução matemática para nos permitir definir “anterior” e “posterior” numa série assim gerada. Na analogia das “frações próprias”, vamos dar o nome “posteridade própria de x com respeito a R ” à classe daqueles termos que pertencem à R -posteridade de algum termo com que x tem a relação R , no sentido que demos antes a “posteridade”, que inclui um termo em sua própria posteridade. Retornando às definições fundamentais, verificamos que “posteridade própria” pode ser definida como se segue:

A “posteridade própria” de x com respeito a R consiste em todos os termos que possuem todas as propriedades hereditárias R possuídas por todos os termos com que x tem a relação R .

Devemos observar que essa definição deve ser formulada de maneira a ser aplicável não só quando há apenas um termo com que x tem a relação R , mas também em casos (como, por exemplo, o de pai e filho) em que pode haver muitos termos com que x tem a relação R . Definimos adicionalmente:

Um termo x é um “ancestral próprio” de y com respeito a R se y pertencer à posteridade própria de x com respeito a R .

Para abreviar, falaremos de “ R -posteridade” e “ R -ancestrais” quando estes termos parecerem mais convenientes.

Retornando agora à geração de séries pela relação R entre termos consecutivos, vemos que, para que esse método seja possível, a relação “ R -ancestral própria” deve ser uma anti-reflexiva, transitiva e conexa. Sob que circunstâncias isso ocorrerá? Ela será sempre transitiva: não importa qual seja a relação R , “ R -ancestral” e “ R -ancestral própria” são sempre ambas transitivas. Mas somente sob certas circunstâncias ela será uma anti-reflexiva ou conexa. Consideremos, por exemplo, a relação de nosso vizinho da esquerda numa mesa de jantar redonda em que há 12 pessoas. Se chamamos isso relação R , a R -posteridade própria de uma pessoa consiste em todas que podem ser alcançadas dando-se a volta da mesa da direita para a esquerda. Isso inclui todas as pessoas à mesa, até mesmo a própria pessoa, uma vez que 12 passos nos trarão de volta a nosso ponto de partida. Assim, num caso como esse, embora a relação “ R -ancestral própria” seja conexa, e embora R seja ela mesma uma anti-reflexiva, não

obtemos uma série porque “ancestral-R própria” não é uma anti-reflexiva. Por essa razão não podemos dizer que uma pessoa vem antes de outra com respeito à relação “à direita de” ou a seu derivado ancestral.

O que vimos foi um exemplo em que a relação ancestral era conexa, mas não contida em diversidade. Um exemplo em que ela é contida em diversidade mas não conexa é derivado do significado usual da palavra “ancestral”. Se x é um ancestral próprio de y , x e y não podem ser a mesma pessoa, mas não é verdade que, de quaisquer duas pessoas, uma deva ser ancestral da outra.

A questão das circunstâncias em que séries podem ser geradas por relações ancestrais derivadas de relações de consecutividade é com frequência importante. Alguns dos casos mais importantes são os seguintes: suponhamos que R-ancestral própria é uma relação, e limitemos nossa atenção à posteridade de algum termo x . Quando assim limitados, a relação “R-ancestral própria” deve ser conexa; portanto, tudo o que resta para assegurar que ela é serial é que seja contida em diversidade. Essa é uma generalização do exemplo da mesa de jantar. Outra generalização consiste em considerar que R é uma relação um-um e inclui tanto a ancestralidade quanto a posteridade de x . Aqui novamente, a única condição necessária para assegurar a geração de uma série é que a relação “R-ancestral própria” seja contida em diversidade.

A geração de ordem por meio de relações de consecutividade, embora importante em sua própria esfera, é menos geral que o método que usa uma relação transitiva para definir a ordem. Acontece frequentemente numa série que haja um número infinito de termos intermediários entre dois que podem ser selecionados, por mais próximos que esses estejam entre si. Tomemos, por exemplo, frações em ordem de magnitude. Entre quaisquer frações há outras — por exemplo, a média aritmética das duas. Conseqüentemente, um par de frações consecutivas é algo que não existe. Se dependêssemos da consecutividade para definir ordem, não seríamos capazes de definir a ordem de magnitude entre frações. Mas, de fato, as relações maior e menor entre frações não exigem geração a partir de relações de consecutividade, e as relações de maior e menor entre frações têm as três características de que precisamos para definir relações seriais. Em todos esses casos a ordem deve ser definida por meio de uma relação *transitiva*, visto que somente uma relação desse tipo é capaz de saltar sobre um número infinito de termos intermediários. O método da

consecutividade, como o de contar para descobrir o número de uma coleção, é apropriado para o finito; pode até ser estendido a certas séries infinitas, a saber, aquelas em que, embora o número total de termos seja infinito, o número de termos entre quaisquer dois é sempre finito; mas não deve ser considerado geral. Não só isso, mas é preciso tomar cuidado para erradicar da imaginação todos os hábitos de pensamento que resultam da suposição de que ele é geral. Se isso não for feito, séries em que não há termos consecutivos permanecerão difíceis e enigmáticas. E séries desse tipo são de importância vital para a compreensão da continuidade, de espaço, tempo e movimento.

Séries podem ser geradas de muitas maneiras, mas tudo depende de se encontrar ou construir uma relação conexa transitiva. Algumas dessas maneiras têm considerável importância. Podemos tomar como ilustrativa a geração de séries por meio de uma relação de três termos que podemos chamar “entre”. Esse método é muito útil em geometria e pode servir como uma introdução para relações que tenham mais que dois termos; a melhor maneira de introduzi-la é em conexão com a geometria elementar.

Dados três pontos quaisquer numa linha reta no espaço comum, deve haver um deles que está *entre* os outros dois. Isso não será o caso com os pontos num círculo ou em qualquer outra curva fechada, porque, dados quaisquer três pontos num círculo, podemos nos deslocar de qualquer um para qualquer outro sem passar pelo terceiro.

De fato, a noção “entre” é característica de séries abertas — ou séries no sentido estrito — em contraposição ao que pode ser chamado de séries “cíclicas”, em que, como no caso das pessoas na mesa de jantar, uma jornada suficiente nos traz de volta ao ponto de partida. Essa noção de “entre” pode ser escolhida como a noção fundamental da geometria comum, mas por enquanto vamos considerar apenas sua aplicação a uma única linha reta e à ordenação de pontos numa linha reta.² Tomando quaisquer dois pontos a , b , a linha (ab) consiste em três partes (afora a e b eles mesmos):

- (1) Pontos entre a e b .
- (2) Pontos x tais que a está entre x e b .
- (3) Pontos y tais que b está entre y e a .

Dessa maneira, a linha (ab) pode ser definida em termos da relação “entre”.

Para que essa relação “entre” possa arranjar os pontos da linha numa ordem da esquerda para a direita, precisamos de certos pressupostos, a saber, os seguintes:

- (1) Se houver alguma coisa entre a e b , a e b não serão idênticos.
- (2) Qualquer coisa que esteja entre a e b estará também entre b e a .
- (3) Qualquer coisa que esteja entre a e b não será idêntica a a (nem, conseqüentemente, a b , em virtude de (2)).
- (4) Se x estiver entre a e b , tudo o que estiver entre a e x estará também entre a e b .
- (5) Se x estiver entre a e b , e b estiver entre x e y , então b estará entre a e y .
- (6) Se x e y estiverem entre a e b , então ou x e y serão idênticos, ou x estará entre a e y , ou x estará entre y e b .
- (7) Se b estiver entre a e x e também entre a e y , então ou x e y serão idênticos, ou x estará entre a e y , ou y estará entre b e x .

Essas sete proposições são obviamente verificadas no caso de pontos numa linha reta em espaço comum. Qualquer relação de três termos que as verifique dará origem a série, com pode ser visto a partir das seguintes definições. No interesse da precisão, suponhamos que a está à esquerda de b . Então os pontos da linha (ab) são (1) aqueles entre os quais b situa-se a — estes vamos chamar de à esquerda de a ; (2) o próprio a ; (3) aqueles entre a e b ; (4) o próprio b ; (5) aqueles entre os quais a situa-se b — estes vamos chamar de à direita de b . Podemos agora definir de maneira geral que de dois pontos x, y , na linha (ab) , diremos que x está “à esquerda de” y em qualquer dos seguintes casos:

- (1) Quando x e y estiverem ambos à esquerda de a , e y estiver entre x e a ;
- (2) Quando x estiver à esquerda de a , e y for a ou b ou estiver entre a e b ou à direita de b ;
- (3) Quando x for a , e y estiver entre a e b ou for b , ou estiver à direita de b ;
- (4) Quando x e y estiverem ambos entre a e b , e y estiver entre x e b ;

- (5) Quando x estiver entre a e b , e y for b ou estiver à direita de b ;
- (6) Quando x for b e y estiver à direita de b ;
- (7) Quando x e y estiverem ambos à direita de b e x estiver entre b e y .

É possível ver que, a partir das sete propriedades que atribuímos à relação “entre”, pode ser deduzido que a relação “à esquerda de”, como definida anteriormente, é uma relação *serial* tal como definimos esse termo. É importante notar que nada nas definições ou no raciocínio depende do que queremos dizer por “entre”, a relação real assim chamada que ocorre no espaço empírico: qualquer relação de três termos que tiver as propriedades puramente formais descritas servirá ao propósito do raciocínio igualmente bem.

A ordem cíclica, como a dos pontos num círculo, não pode ser gerada por meio de relações de três termos de “entre”. Precisamos de uma relação de quatro termos, que pode ser chamada “separação de pares”. A idéia pode ser ilustrada considerando-se uma viagem de volta ao mundo. Podemos ir da Inglaterra para a Nova Zelândia passando por Suez ou passando por São Francisco; não podemos dizer com precisão que nenhum desses dois lugares está “entre” a Inglaterra e a Nova Zelândia. Mas se um homem escolhe essa rota para dar a volta ao mundo, seja qual for o caminho que tome, os momentos que passa na Inglaterra e na Nova Zelândia são separados uns dos outros pelos momentos que passa por Suez ou por São Francisco, e inversamente. Generalizando, se tomarmos quaisquer quatro pontos num círculo, poderemos separá-los em dois pares, digamos a e b e x e y , tais que, para ir de x para y se deve passar por a ou por b . Sob essas circunstâncias, dizemos que o par (a, b) está “separado” pelo par (x, y) . A partir dessa relação pode ser gerada uma ordem cíclica, de uma maneira que se assemelha àquela pela qual geramos uma ordem aberta a partir de “entre”, porém um pouco mais complicada.³

O objetivo da segunda metade deste capítulo foi sugerir o tema que podemos chamar “geração de relações seriais”. Quando tais relações foram definidas, a geração delas a partir de outras relações que possuam apenas algumas das propriedades requeridas para séries torna-se muito importante, especialmente na filosofia da geometria e da física. Mas não podemos, nos limites do presente volume, fazer mais do que tornar o leitor ciente da existência desse tema.

Capítulo 5

Tipos de relação

Grande parte da filosofia da matemática diz respeito a relações, e muitos diferentes tipos de relações têm diferentes tipos de usos. Ocorre com frequência que uma propriedade que pertence a *todas* as relações só seja importante com respeito a relações de certas espécies; nesses casos o leitor não verá a relevância da proposição que afirma tal propriedade, a menos que tenha em mente as espécies de relações para as quais ela é útil. No interesse dessa descrição, bem como em razão do interesse intrínseco do tema, convém termos em mente uma lista elementar das variedades de relações mais úteis matematicamente.

Lidamos no capítulo anterior com uma classe sumamente importante, a saber, relações *seriais*. Cada uma das três propriedades que combinamos ao definir série — a saber, *assimetria*, *transitividade* e *conexidade* — tem sua própria importância. Começaremos dizendo alguma coisa sobre cada uma dessas três.

Assimetria, isto é, a propriedade de ser incompatível com o inverso, é uma característica do maior interesse e importância. Para desenvolver suas funções, vamos considerar vários exemplos. A relação *marido* é assimétrica, e assim também a relação *esposa*; isto é, se a é marido de b , b não pode ser marido de a , e similarmente no caso de esposa. Por outro lado, a relação “cônjuge” é simétrica: se a é cônjuge de b , então b é cônjuge de a . Suponhamos agora que nos foi dada a relação *cônjuge* e desejamos derivar a relação *marido*. *Marido* é o mesmo que *cônjuge do sexo masculino* ou *cônjuge de uma mulher*; assim, a relação *marido* pode ser derivada de *cônjuge*, ou limitando-se o domínio a homens ou limitando-se o inverso a mulheres. Vemos a partir desse exemplo que, quando uma relação simétrica é dada, por vezes é possível, sem a ajuda de nenhuma relação adicional, separá-la em duas relações assimétricas. Mas os casos em que isso é possível são raros e excepcionais: são casos em que

há duas classes mutuamente exclusivas, digamos α e β , tais que, seja qual for a relação existente entre dois termos, um dos termos é membro de α e o outro é membro de β — como, no caso de *cônjuge*, um termo da relação pertence à classe dos homens, e um, à classe das mulheres. Nesse caso, a relação com seu domínio limitado a α será assimétrica, e o mesmo ocorrerá com a relação com seu domínio limitado a β . Mas não são casos dessa espécie que ocorrem quando estamos lidando com séries de mais que dois termos; pois, numa série, todos os termos, exceto o primeiro e o último (se eles existirem), pertencem tanto ao domínio quanto ao domínio inverso da relação geradora, de modo que uma relação como *marido*, em que o domínio e do domínio inverso não se superpõem, está excluída.

A questão de como *construir* relações que tenham alguma propriedade útil por meio de operações com relações que possuem apenas rudimentos da propriedade é de considerável importância. É fácil construir transitividade e conexidade em muitos casos em que a relação originalmente dada não as possui: por exemplo, se R for uma relação qualquer, a relação ancestral derivada de R por indução generalizada será transitiva; e se R for uma relação muitos-um, a relação ancestral será conexa se estiver limitada à posteridade de um dado termo. No entanto, obter a propriedade da assimetria por construção é muito mais difícil. O método pelo qual derivamos *marido* e *cônjuge*, como vimos, não está disponível nos casos mais importantes, como *maior*, *antes*, *à direita de*, em que domínio e domínio inverso se superpõem. Em todos esses casos, podemos é claro obter uma relação assimétrica somando a relação dada à sua inversa, mas não podemos passar de volta dessa relação simétrica para a relação simétrica original, exceto com a ajuda de alguma relação assimétrica. Tomemos, por exemplo, a relação *maior*: a relação *maior ou menor* — isto é, *desigual* — é simétrica, mas não há nada nela que mostre que é a soma de duas relações assimétricas. Tomemos uma relação como “diferente no formato”. Isso não é a soma de uma relação assimétrica e seu inverso, uma vez que formatos não formam uma série única; mas não há nada para mostrar que ela difere de “diferente na magnitude” se já não soubéssemos que magnitudes têm relações de maior e menor. Isso ilustra o caráter fundamental da assimetria como uma propriedade das relações.

Do ponto de vista da classificação das relações, ser assimétrico é uma característica muito mais importante do que implicar diversidade. Relações assimétricas implicam diversidade, mas o contrário não se

verifica. “Desigual”, por exemplo, implica diversidade, mas é simétrico. De maneira geral, podemos dizer que, se desejássemos prescindir na medida do possível das proposições relacionais e substituí-las por aquelas que atribuem predicados a sujeitos, conseguiríamos fazê-lo na medida em que nos limitássemos a relações *simétricas*: aquelas que não implicam diversidade, se forem transitivas, podem ser vistas como afirmando um predicado comum, ao passo que aquelas que implicam diversidade podem ser vistas como afirmando predicados incompatíveis. Por exemplo, considere a relação de *similaridade entre classes*, por meio da qual definimos números. Essa relação é simétrica, é transitiva e não implica diversidade. Seria possível, embora menos simples que o procedimento que adotamos, ver o número de uma coleção como um predicado da coleção: nesse caso duas classes similares serão duas classes que tenham o mesmo predicado numérico, ao passo que duas classes não similares serão duas classes com predicados numéricos diferentes. Tal método de substituir relações por predicados é formalmente possível (embora com freqüência muito inconveniente), contanto que as relações envolvidas sejam simétricas; mas é formalmente impossível quando elas são assimétricas, porque tanto a igualdade quanto a diferença de predicados são simétricas. As relações assimétricas são, podemos dizer, as mais caracteristicamente relacionais das relações, e o mais importante para o filósofo que deseja estudar a natureza lógica fundamental das relações.

Outra classe de relações da maior utilidade é a classe das relações um-muitos, isto é, relações que no máximo um termo pode ter com dado termo. Assim são pai, mãe, marido (exceto no Tibete), quadrado de, seno de, e assim por diante. Mas genitor, raiz quadrada, e assim por diante, não são um-muitos. É possível, formalmente, substituir todas as relações por relações um-muitos por meio de um truque. Tomemos (digamos) a relação *menor* entre os números indutivos. Dado qualquer número n maior do que 1, não haverá apenas um número que tenha a relação menor do que n , mas podemos formar toda a classe de números menores do que n . Essa é uma classe, e sua relação com n não é partilhada por nenhuma outra. Podemos chamar a classe dos números menores do que n a “ancestralidade própria” de n , no sentido em que falamos de ancestralidade e posteridade em conexão com a indução matemática. Assim, “ancestralidade própria” é uma relação um-muitos (*um-muitos* será sempre usado de modo a incluir *um-um*), visto que cada número determina uma única classe de números

como constituindo sua ancestralidade própria. Assim, a relação *menor que* pode ser substituída por *ser membro da ancestralidade própria de*. Dessa maneira, uma relação um-muitos em que o um é uma classe, juntamente com a qualidade de membro dessa classe, pode sempre substituir formalmente uma relação que não é um-muitos. Peano, que, por alguma razão sempre concebe intuitivamente uma relação como um-muitos, trata dessa maneira aquelas que não são naturalmente assim. A redução a relações um-muitos por esse método, no entanto, embora possível como matéria de forma, não representa uma simplificação técnica, e temos todas as razões para pensar que não representa uma análise filosófica, ainda que apenas porque classes devem ser vistas como “ficções lógicas”. Continuaremos, portanto, a considerar relações um-muitos como um tipo especial de relações.

Relações um-muitos estão envolvidas em todas as expressões da forma “isso e aquilo disso e daquilo”. “O rei da Inglaterra” ou “a mulher de Sócrates”, “o pai de John Stuart Mill”, e assim por diante, descrevem todas alguma pessoa por meio de uma relação um-muitos com um dado termo. Uma pessoa não pode ter mais que um pai, portanto “o pai de John Stuart Mill” descrevia uma determinada pessoa, mesmo que não soubéssemos quem era. Há muito a dizer sobre o tema das descrições, mas por enquanto são as relações que nos interessam, e as descrições são relevantes apenas enquanto exemplificam os usos de relações um-muitos. Convém observar que todas as funções matemáticas resultam de relações um-muitos: o logaritmo de x , o co-seno de x etc., são, como o pai de x , termos descritos por meio de uma relação um-muitos (logaritmo, co-seno etc.) com um termo dado (x). A noção de *função* não precisa ser limitada a números ou aos usos a que os matemáticos nos acostumaram; pode ser estendida a todos os casos de relações um-muitos, e “o pai de x ” é uma função de que x é o argumento de maneira tão legítima quanto “o logaritmo de x ”. Nesse sentido, funções são funções *descritivas*. Como veremos mais adiante, há funções de uma espécie ainda mais geral e mais fundamental, a saber, funções *proposicionais*; mas por enquanto limitaremos nossa atenção a funções descritivas, isto é, “o termo que tem a relação R com x ”, ou, para abreviar, “a R de x ”, onde R é qualquer relação um-muitos.

Deve-se observar que para que “a R de x ” descreva um termo preciso, x deve ser um termo com que alguma coisa tem a relação R, e não deve

haver mais que um termo tendo a relação R com x , uma vez que “o”, corretamente utilizado, deve implicar singularidade. Assim podemos falar de “o pai de x ” se x for algum ser humano exceto Adão e Eva; mas não podemos falar de “o pai de x ” se x for uma mesa ou uma cadeira ou qualquer outra coisa que não tem um pai. Diremos que a R de x “existe” quando há apenas um termo, e não mais, que tem a relação R com x . Assim se R for uma relação um-muitos, a R de x existe sempre que x pertencer ao domínio inverso de R , e não de outro modo. No tocante a “a R de x ” como uma função no sentido matemático, dizemos que x é o “argumento” da função, e se y for o termo que tem a relação com x , isto é, se y for a R de x , então y será o “valor” da função para o argumento x . Se R for uma relação um-muitos, o âmbito de argumentos possíveis para a função será o domínio inverso de R , e o âmbito de valores será o domínio. Assim, o âmbito de argumentos possíveis para a função “o pai de x ” são todos os que têm pais, isto é, o domínio inverso da relação *pai*, ao passo que o âmbito de valores possíveis para a função é todos os pais, isto é, o domínio da relação.

Muitas das noções mais importantes na lógica das relações são funções descritivas, por exemplo: inverso, domínio, domínio inverso, campo. Outros exemplos ocorrerão à medida que prosseguirmos.

Entre as relações um-muitos, as relações *um-um* são uma classe especialmente importante. Já tivemos oportunidade de falar de relações um-um em conexão com a definição de número, mas é necessário ter familiaridade com elas, e não meramente conhecer sua definição formal. Sua definição formal pode ser derivada da definição de relações um-muitos: elas podem ser definidas como relações um-muitos que são também o inverso das relações um-muitos, isto é, como relações que são tanto um-muitos como muitos-um. Relações um-muitos podem ser definidas como aquelas em que, se x tiver a relação em questão com y , não há nenhum outro termo x' que também tenha a relação com y . Podem também ser definidas da seguinte maneira: dados dois termos x e x' , os termos com que x tem a relação dada e aqueles com que x' a tem não têm nenhum membro em comum. Ou ainda, podem ser definidas como relações tais que o produto relativo de uma delas e seu inverso implica identidade, onde o “produto relativo” de duas relações R e S é aquela relação existente entre x e z quando há um termo intermediário y , de tal modo que x tenha a relação R com y e y tenha a relação S com z . Assim,

por exemplo, se R for a relação de pai com filho, o produto relativo de R e seu inverso será a relação existente entre x e um homem z quando há uma pessoa y , tal que x seja o pai de y e y seja o filho de z . É óbvio que x e z devem ser a mesma pessoa. Se, por outro lado, tomarmos a relação entre genitor e filho, que não é um-muitos, não podemos mais afirmar que, se x for pai de y e y for filho de z , x e z devem ser a mesma pessoa, porque um pode ser o pai de y e o outro, a mãe. Isso ilustra o que é característico das relações um-muitos quando o produto de uma relação e seu inverso implica identidade. No caso de relações um-um isso acontece, e também o produto relativo do inverso e a relação implica identidade. Dada uma relação R , é conveniente, se x tiver a relação R com y , conceber y como sendo alcançado a partir de x por um “ R -passo” ou um “ R -vetor”. No mesmo caso, x será alcançado a partir de y por um “passo R para trás”. Podemos, portanto, expressar a característica das relações um-muitos de que estivemos tratando dizendo que um passo R seguido por um passo R para trás deve nos levar ao nosso ponto de partida. Com outras relações, isso não ocorre em absoluto; por exemplo, se R for a relação de um filho com um genitor, o produto relativo de R e seu inverso será a relação “a própria pessoa ou irmão ou irmã”, e se for a relação de neto com avô, o produto relativo de R e seu inverso será a relação “a própria pessoa ou irmão ou irmã ou primo em primeiro grau”. Convém observar que o produto relativo de duas relações em geral não é comutativo, isto é, o produto relativo de R e S não é em geral a mesma relação que o produto de S e R . Por exemplo, o produto relativo de um genitor e irmão é tio, mas o produto relativo de irmão e genitor é genitor.

Relações um-um dão uma correlação de duas classes, termo por termo, de tal modo que cada termo em cada uma das classes tem seu correlato na outra. É mais simples compreender essas correlações quando as duas classes não têm membros em comum, como a classe dos maridos e a classe das esposas, pois, nesse caso, sabemos de imediato se um termo deve ser considerado um termo *do qual* a relação de correlação vem, ou um termo *para* o qual ela vai. Assim se x e y forem marido e mulher, então, com respeito à relação “marido” x é referente e y é referido, mas com respeito à relação “esposa”, y é referente, e x , referido. Dizemos que uma relação e seu inverso e têm “sentidos” opostos; assim o “sentido” de uma relação que vai de x para y é o inverso daquele da relação correspondente de y para x . O fato de uma relação ter um “sentido” é

fundamental, e é parte da razão por que a ordem pode ser gerada por relações adequadas. Convém observar que a classe de todos os referentes possíveis para a uma dada relação é seu domínio, e a classe de todos os referidos possíveis é seu domínio inverso.

Mas acontece com muita frequência que o domínio e o domínio inverso de uma relação um-um se superponham. Tomemos, por exemplo, os primeiros dez números inteiros (excluindo 0) e acrescentemos 1 a cada um; assim, em vez dos dez primeiros números inteiros temos agora os números inteiros

2, 3, 4, 5, 6, 7, 8, 9, 10, 11.

Esses são os mesmos que tínhamos antes, exceto que 1 foi cortado no início e 11 foi acrescentado no fim. Ainda há dez números inteiros: eles estão correlacionados aos dez anteriores pela relação de n com $n + 1$, que é uma relação um-um. Poderíamos igualmente, em vez de acrescentar 1 a cada um de nossos dez números inteiros originais, ter dobrado cada um deles, obtendo assim os números inteiros

2, 4, 6, 8, 10, 12, 14, 16, 18, 20.

Ainda temos aqui cinco números de nosso conjunto anterior de números inteiros, a saber, 2, 4, 6, 8, 10. A relação de correlação nesse caso é a relação de um número com seu duplo, que é novamente uma relação um-um. Ou poderíamos ter substituído cada número por seu quadrado, obtendo assim o conjunto

1, 4, 9, 16, 25, 36, 49, 64, 81, 100

Dessa vez restam apenas três números de nosso conjunto inicial, a saber, 1, 4, 9. Processos de correlação como esses podem ser variados interminavelmente. O caso mais interessante do tipo descrito é aquele em que nossa relação um-um tem um domínio inverso que é parte, mas não o todo, do domínio. Se, em vez de limitar o domínio aos dez primeiros números inteiros, tivéssemos considerado a totalidade dos números indutivos, os exemplos mencionados teriam ilustrado esse caso. Podemos dispor os números envolvidos em duas fileiras, pondo o correlato

diretamente embaixo do número do qual ele é correlato. Assim quando o correlator for a relação de n com $n + 1$, teremos duas fileiras:

$$\begin{array}{c} 1, 2, 3, 4, 5, \dots n \dots \\ 2, 3, 4, 5, 6, \dots n + 1 \dots \end{array}$$

Quando o correlator for a relação de um número com seu duplo, teremos as duas colunas:

$$\begin{array}{c} 1, 2, 3, 4, 5, \dots n \dots \\ 2, 4, 6, 8, 10, \dots 2n \dots \end{array}$$

Quando o correlator for a relação de um número com seu quadrado, as fileiras serão:

$$\begin{array}{c} 1, 2, 3, 4, 5, \dots n \dots \\ 1, 4, 9, 16, 25, \dots n^2 \dots \end{array}$$

Em todos esses casos, todos os números indutivos ocorrem na fileira de cima, e apenas alguns na fileira de baixo.

Casos dessa espécie, em que o domínio inverso é uma “parte própria” do domínio (isto é, uma parte, não a totalidade), voltarão a nos ocupar quando chegarmos a tratar da infinidade. Por enquanto, desejamos apenas observar que eles existem e requerem consideração.

Uma outra classe de correlações muitas vezes importantes é a classe chamada “permutações”, em que o domínio e o domínio inverso são idênticos. Consideremos, por exemplo, os seis arranjos possíveis das três letras:

a, b, c
a, c, b
b, c, a
b, a, c
c, a, b
c, b, a

Cada um desses arranjos pode ser obtido a partir de qualquer dos outros por meio de uma correlação. Tomemos, por exemplo, o primeiro e o

último: (a, b, c) e (c, b, a) . Aqui a está correlacionado a c , b consigo mesmo e c a a . É óbvio que a combinação de duas permutações é novamente uma permutação, isto é, a permutação de uma dada classe forma o que é chamado um “grupo”.

Esses vários tipos de correlações têm importância em várias conexões, algumas para um objetivo, algumas para outro. A noção geral de correlações um-um tem inesgotável importância na filosofia da matemática, como já vimos em parte, mas veremos de maneira muito mais completa à medida que avançarmos. Um de seus usos nos ocupará no próximo capítulo.

Capítulo 6

Similaridade das relações

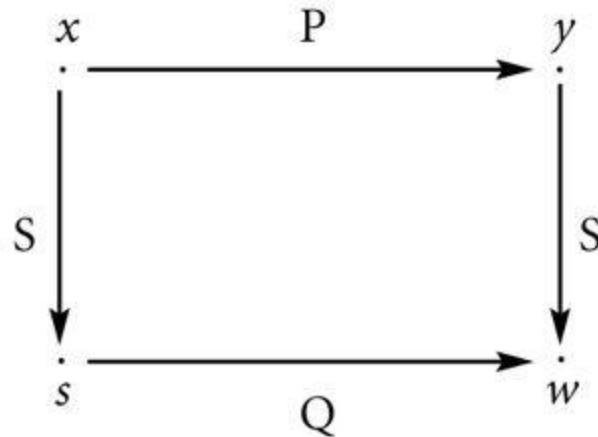
Vimos no Capítulo 2 que duas classes têm o mesmo número de termos quando são “similares”, isto é, quando há uma relação um-um cujo domínio é uma classe e cujo domínio inverso é a outra. Num caso assim podemos dizer que há uma “correlação um-um” entre as duas classes.

Nesse capítulo temos de definir uma relação entre relações, que desempenhará para elas o mesmo papel que a similaridade desempenha para classes. Chamaremos isso “similaridade de relações”, ou “semelhança” quando parecer desejável usar uma palavra diferente da que usamos para classes. Como devemos definir semelhança?

Continuaremos empregando a noção de correlação: vamos dar por certo que o domínio de uma das relações pode ser correlacionado com o domínio da outra, e o domínio inverso com o domínio inverso; mas isso não é suficiente para a espécie de semelhança que desejamos ter entre nossas duas relações. O que desejamos é que, sempre uma das duas relações exista entre dois termos, a outra exista entre os correlatos desses dois termos. O exemplo mais fácil do tipo de coisa que desejamos é um mapa. Quando um lugar está ao norte de outro, o lugar no mapa correspondente a um está acima do lugar no mapa correspondente ao outro; quando um lugar está a oeste de outro, o lugar no mapa correspondente a um está à esquerda do lugar no mapa correspondente ao outro e assim por diante. A estrutura do mapa corresponde à do país que ele representa. As relações espaciais no mapa têm uma “semelhança” com as relações espaciais no país mapeado. Esse é o tipo de conexão entre relações que desejamos definir.

Podemos, em primeiro lugar, introduzir vantajosamente uma certa restrição. Vamos nos limitar, ao definir semelhança, a relações que tenham “campos”, isto é, aquelas que permitam a formação de uma única classe a partir do domínio e do domínio inverso. Esse nem sempre é o caso.

Tomemos, por exemplo, a relação “domínio”, isto é, a relação que o domínio de uma relação tem com ela. Essa relação tem todas as classes como seu domínio, uma vez que toda classe é o domínio de alguma relação; e tem todas as relações como seu domínio inverso, uma vez que toda relação tem um domínio. Mas classes e relações não podem ser somadas para formar uma nova classe única, porque são de “tipos” lógicos diferentes. Não precisamos entrar na difícil doutrina dos tipos, mas é bom saber quando nos abstermos de nos aventurar nela. Podemos dizer, sem considerar os fundamentos para a asserção, que uma relação só tem “um campo” quando é o que chamamos “homogênea”, isto é, quando seu domínio e seu domínio inverso são do mesmo tipo lógico; e como uma indicação grosseira do que entendemos por um “tipo”, podemos dizer que indivíduos, classes de indivíduos, relações entre indivíduos, relações entre classes, relações de classes com indivíduos, e assim por diante, são todos diferentes tipos. Mas a noção de semelhança não é muito útil quando aplicada a relações não homogêneas; por isso, ao definir semelhança, simplificaremos nosso problema falando do “campo” de uma das relações envolvidas. Isso limita em certa medida a generalidade de nossa definição, mas a limitação não tem nenhuma importância prática, e depois de ter sido formulada, não precisa mais ser lembrada. Podemos definir duas relações P e Q como “similares”, ou como tendo “semelhança”, quando há uma relação S um-um cujo domínio é o campo de P e cujo domínio inverso é o campo de Q , e que é tal que, se um termo tiver a relação P com outro, o correlato dele terá a relação Q com o correlato do outro, e vice-versa. Uma figura tornará isto mais claro. Suponhamos que x e y são dois termos com a relação P . Deve haver, portanto, dois termos z , w , tais que x tem a relação S com z , y tem a relação S com w e z tem a relação Q com w . Se isso acontecer com cada par de termos como z e w , é claro que para cada caso em que existir a relação P existirá um caso correspondente da relação Q , e vice-versa; e isso é o que desejamos assegurar com nossa definição. Podemos eliminar algumas redundâncias no esboço de definição citado, observando que, quando as condições mencionadas se realizam, a relação P é igual ao produto relativo de S e Q e o inverso de S , isto é, o passo P de x para y pode ser substituído pela sucessão do passo S de x para z , o passo Q de z para w , e o passo atrás S de w para y . Podemos assim estabelecer as seguintes definições.



Diz-se que uma relação S é um “correlator” ou um “correlator ordinal” de duas relações P e Q se S for um-um, tiver o campo de Q por seu domínio inverso, e for tal que P é o produto relativo de S e Q e o inverso de S . Diz-se que duas relações P e Q são “similares” ou têm “semelhança” quando há no mínimo um correlator P e Q .

Veremos que essas definições proporcionam o que estipulamos anteriormente ser necessário.

Veremos que, quando duas relações são similares, elas partilham todas as propriedades que não dependem dos termos reais em seus campos. Por exemplo, se uma implica diversidade, a outra também o faz; se uma é transitiva, a outra também é; se uma é conexa, a outra também. Em consequência, se uma é serial, a outra também é. Mais uma vez, se uma é um-muitos ou um-um, a outra é um-muitos ou um-um; e assim por diante, com todas as propriedades gerais das relações. Mesmo afirmações que envolvam os termos reais do campo de uma relação, ainda que possam não ser verdadeiras tal como são quando aplicadas a uma relação similar, serão sempre passíveis de tradução em afirmações análogas. Essas considerações nos conduzem a um problema que tem, em filosofia matemática, uma importância que até agora não foi em absoluto adequadamente reconhecida. Nosso problema pode ser formulado da seguinte maneira: dada alguma afirmação numa língua da qual conhecemos a gramática e a sintaxe, mas não o vocabulário, quais são os significados possíveis dessa afirmação, e quais são os significados das palavras desconhecidas que a tornariam verdadeira?

Essa questão é importante porque representa, muito mais fielmente do que se poderia supor, o estado de nosso conhecimento da natureza.

Sabemos que certas proposições científicas — que, nas ciências mais avançadas, são expressas em símbolos matemáticos — são mais ou menos verdadeiras em relação ao mundo, mas estamos extremamente confusos quanto à interpretação a ser dada aos termos que ocorrem nessas proposições. Sabemos muito mais (para usar, por um momento, um par de termos antiquado) sobre a *forma* do que sobre a *matéria* da natureza. Assim, o que realmente sabemos quando enunciamos uma lei da natureza é apenas que há provavelmente *alguma* interpretação de nossos termos que tornará a lei aproximadamente verdadeira. A grande importância prende-se à questão: quais são os significados possíveis de uma lei expressa em termos de que não conhecemos o significado substantivo, mas apenas a gramática e a sintaxe? E é essa a questão sugerida previamente.

Por enquanto, vamos ignorar a questão geral, que voltará a nos ocupar num estágio posterior; o tema da semelhança precisa ser ele próprio mais investigado.

Em razão do fato de que, quando duas relações são similares, suas propriedades são as mesmas exceto quando dependem de os campos serem compostos exatamente dos termos de que elas são compostas, é desejável ter uma nomenclatura que reúna todas as relações similares a uma relação dada. Assim como chamamos o conjunto daquelas classes similares a uma dada classe o “número” dessa classe, assim também podemos chamar o conjunto de todas as relações similares a uma relação dada o “número” dessa relação. Para evitar confusão com os números apropriados a classes, porém, falaremos neste caso de um “número de relação”. Temos assim as seguintes definições: o “número de relação” de uma dada relação é a classe de todas as relações similares à relação dada; e “números de relação” são o conjunto de todas aquelas classes de relações que são números de relação de várias relações, ou, o que dá no mesmo, um número de relação é uma classe de relações que consiste em todas aquelas relações que são similares a um membro da classe.

Quando é necessário falar dos números de classes de uma maneira que torne impossível confundi-los com números de relação, nós os chamaremos de “números cardinais”. Números cardinais são portanto os números apropriados a classes. Esses incluem os números inteiros comuns da vida cotidiana, e também certos números infinitos, de que falaremos mais tarde. Quando falamos de “números” sem qualificação, deve-se compreender que estamos nos referindo a números cardinais. A definição

de um número cardinal, convém lembrar, é a seguinte: o “número cardinal” de uma dada classe é o conjunto de todas aquelas classes similares à classe dada.

A aplicação mais óbvia de números de relação é a *séries*. Duas séries podem ser consideradas igualmente longas quando têm o mesmo número de relação. Duas séries *finitas* terão o mesmo número de relação quando seus campos tiverem o mesmo número cardinal de termos, e somente então — isto é, uma série de (digamos) 15 termos terá o mesmo número de relação que qualquer outra série de 15 termos, mas não terá o mesmo número de relação que uma série de 14 ou 16 termos, nem, é claro, o mesmo número de relação que uma relação que não seja serial. Assim, no caso muito especial das séries finitas, há paralelismo entre números cardinais e números de relação. Os números de relação aplicáveis a séries podem ser chamados “números seriais” (os comumente chamados “números ordinais” são uma subclasse desses); assim, um número serial finito é determinado quando conhecemos o número cardinal de termos no campo de uma série que tenha o número serial em questão. Se n for um número cardinal finito, o número de relação de uma série que tem n termos será chamado o número “ordinal” n . (Há também números ordinais infinitos, mas falaremos deles em um capítulo posterior.) Quando o número cardinal de termos no campo de uma série for infinito, o número de relação da série não será determinado meramente pelo número cardinal; de fato, existe um número infinito de números de relação para um número cardinal infinito, como veremos quando passarmos a considerar séries infinitas. Quando uma série é infinita, o que podemos chamar de seu “comprimento”, isto é, seu número de relação, pode variar sem mudança no número cardinal; mas quando uma série é finita, isso não pode acontecer.

Podemos definir adição e multiplicação para números de relação, bem como para números cardinais, e toda uma aritmética de números de relação pode ser desenvolvida. É fácil ver como isso deve ser feito considerando o caso das séries. Suponhamos, por exemplo, que desejamos definir a soma de duas séries que não se superpõem de tal maneira que o número de relação da soma possa ser definido como a soma dos números de relação das duas séries. Em primeiro lugar, é claro que há uma *ordem* envolvida como entre as duas séries: uma delas deve ser posta antes da outra. Assim se P e Q forem as relações geradoras das duas séries, na série

que é a soma delas com P posto antes de Q , todos os membros do campo de P vão preceder todos os membros do campo de Q . Dessa maneira, a relação serial que deve ser definida como a soma de P e Q não é simplesmente “ P ou Q ”, mas “ P ou Q ou a relação de qualquer membro do campo de P com qualquer membro do campo de Q ”. Supondo que P e Q não se superpõem, essa relação é serial, mas “ P ou Q ” não é serial, não sendo conexa, uma vez que não existe entre um membro do campo de P e um membro do campo de Q . Portanto, a soma de P e Q , tal como definida, é aquilo de que precisamos para definir a soma de dois números de relação. Modificações similares são necessárias para produtos e potências. A aritmética resultante não obedece à lei comutativa: a soma ou produto de dois números de relação depende em geral da ordem em que são tomados. Mas ela obedece à lei associativa, uma forma da lei distributiva, e a duas das leis formais para as potências, não só como aplicadas a números seriais, mas como aplicadas a números de relação em geral. De fato, embora recente, a aritmética das relações é um ramo perfeitamente respeitável da matemática.

Não se deve supor, meramente porque séries fornecem a aplicação mais óbvia da idéia de semelhança, que não há outras aplicações importantes. Já mencionamos mapas, e poderíamos estender nossos pensamentos dessa ilustração para a geometria em geral. Se o sistema de relações pelo qual a geometria é aplicada a certo conjunto de termos puder ser transposto inteiramente para relações de semelhança com um sistema aplicando-se a outro conjunto de termos, então a geometria dos dois conjuntos será indistinguível do ponto de vista matemático, isto é, as proposições serão as mesmas, exceto quanto ao fato de que serão aplicadas em um caso a um conjunto de termos e no outro a um outro. Podemos ilustrar isso pelas relações da espécie que pode ser chamada “entre”, que consideramos no Capítulo 4. Vimos ali que uma relação de três termos, contanto que tenha certas propriedades lógicas formais, dará origem a série, e pode ser chamada uma “relação entre”. Dados dois pontos quaisquer, podemos usar a relação-entre para definir a linha reta determinada por aqueles dois pontos; ela consiste em a e b juntamente com todos os pontos x , tais que a relação-entre exista entre os três pontos a , b , x em uma ordem ou outra. O. Veblen demonstrou que podemos considerar todo o nosso espaço como o campo de uma relação entre de três termos, e definir nossa geometria pelas propriedades que atribuímos a nossas relações-entre.¹ Ora, é tão fácil

definir a semelhança entre relações de três termos quanto em relações de dois. Se B e B' forem duas relações-entre, de modo que " $xB(y, z)$ " significa " x está entre y e z com respeito a B ", poderemos chamar S um correlator de B e B' se ele tiver o campo de B' por seu domínio inverso, e for tal que a relação B exista entre três termos quando B' existir entre seus correlatos S e somente então. E diremos que B é semelhante a B' quando houver pelo menos um correlator de B com B' . O leitor pode se convencer facilmente de que, se B for semelhante a B' nesse sentido, não poderá haver nenhuma diferença entre a geometria gerada por B e aquela gerada por B' .

Disso se segue que o matemático não precisa se preocupar com o ser particular ou a natureza intrínseca de seus pontos, linhas e planos, mesmo quando está especulando como um matemático *aplicado*. Podemos dizer que há uma evidência empírica da verdade aproximada daquelas partes da geometria que não são matéria de definição. Mas não há nenhuma evidência empírica quanto ao que deve ser um "ponto". Ele tem de ser algo que satisfaça tão plenamente quanto possível nossos axiomas, mas não precisa ser "muito pequeno" ou "sem partes", isso é indiferente. Se pudermos, a partir de material empírico, construir uma estrutura lógica, não importa qual seja seu grau de complicação, que satisfaça nossos axiomas geométricos, essa estrutura poderá ser legitimamente chamada um "ponto". Não devemos dizer que não há nada mais que poderia ser legitimamente chamado um "ponto"; devemos dizer apenas: "Esse objeto que construímos é suficiente para o geômetra; pode ser um de muitos objetos, quaisquer dos quais seria suficiente, mas isso não nos preocupa, visto que esse objeto é suficiente para justificar a verdade empírica da geometria, na medida em que a geometria não é matéria de definição." Isso é apenas uma ilustração do princípio geral segundo o qual o que importa na matemática, e numa medida muito grande na ciência física, não é a natureza intrínseca de nossos termos, mas a natureza lógicas de suas inter-relações.

Podemos dizer, sobre duas relações similares, que elas têm a mesma "estrutura". Para propósitos matemáticos (embora não para os da filosofia pura), a única coisa importante acerca de uma relação são os casos em que ela existe, não sua natureza intrínseca. Assim como uma classe pode ser definida por vários conceitos diferentes mas co-extensivos — por exemplo, "homem" e "bípede implume" —, assim também duas relações

conceitualmente diferentes podem existir no mesmo conjunto de casos. Um “caso” em que uma relação existe deve ser concebido como um par de termos, com uma ordem, de tal modo que um dos termos venha em primeiro lugar e o outro em segundo; o par deve, é claro, ser tal que seu primeiro termo tenha a relação em questão com seu segundo. Tomemos (digamos) a relação “pai”: é possível definir o que podemos chamar a “extensão” dessa relação como a classe todos os pares ordenados (x, y) que são tais que x é o pai de y . Do ponto de vista matemático, a única coisa de importância na relação “pai” é que ela define esse conjunto de pares ordenados. Em geral, dizemos: a “extensão” de uma relação é a classe daqueles pares ordenados (x, y) que são tais que x tem a relação em questão com y .

Podemos agora dar mais um passo no processo de abstração, e considerar o que entendemos por “estrutura”. Dada qualquer relação, se ela for suficientemente simples, podemos construir um mapa dela. No interesse da precisão, tomemos uma relação cuja extensão sejam os seguintes pares: $ab, ac, ad, bc, ce, dc, de$, onde a, b, c, d, e são cinco termos, não importa quais. Podemos fazer um “mapa” dessa relação tomando cinco pontos num plano e conectando-os com setas, como na figura ao lado. O que o mapa revela é o que chamamos a “estrutura” da relação.

entra fig. da pág. 60 (OI)

É claro que a “estrutura” da relação não depende dos termos particulares que compõem o campo da relação. O campo pode ser alterado sem que a estrutura se altere, e a estrutura pode ser alterada sem alteração no campo — por exemplo, se acrescentássemos o par ae à ilustração acima, alteraríamos a estrutura, mas não o campo. Duas relações têm a mesma “estrutura”, diremos, quando o mesmo mapa servir para ambas — ou, o que dá no mesmo, quando cada uma pode ser um mapa para a outra (já que toda relação pode ser seu próprio mapa). E isso, como um momento de reflexão mostra, é exatamente a mesma coisa que chamamos “semelhança”. Quer dizer, duas relações têm a mesma estrutura quando têm semelhança, isto é, quando têm o mesmo número de relação. Assim, o que definimos como “número de relação” é exatamente a mesma coisa que é obscuramente significada pela palavra “estrutura” — uma palavra que,

embora importante, nunca (ao que saibamos) é definida em termos precisos por aqueles que a utilizam.

Muita especulação na filosofia tradicional poderia ter sido evitada se a importância da estrutura e a dificuldade de penetrá-la tivessem sido compreendidas. Por exemplo, muitas vezes se diz que espaço e tempo são subjetivos, mas eles têm correspondentes objetivos; ou que fenômenos são subjetivos, mas são causados pelas coisas em si mesmas, que devem ter diferenças *inter se* correspondentes às diferenças nos fenômenos a que dão origem. Quando tais hipóteses são feitas, supõe-se em geral que podemos saber muito pouco sobre os correspondentes objetivos. Na realidade, contudo, se as hipóteses tal como formuladas estivessem corretas, os correspondentes objetivos formariam um mundo dotado da mesma estrutura que o mundo fenomenal, e nos permitiriam inferir dos fenômenos a verdade de todas as proposições que podem ser formuladas em termos abstratos e são conhecidas como expressando a verdade dos fenômenos. Se o mundo fenomenal tem três dimensões, o mesmo deve ocorrer com o mundo por trás dos fenômenos; se o mundo fenomenal é euclidiano, assim também deve ser o outro; e assim por diante. Em suma, toda proposição dotada de uma significação comunicável deve ser verdadeira acerca de ambos os mundos ou de nenhum deles: a única diferença deve residir exatamente nessa essência de individualidade que sempre escapa às palavras e frustra a descrição, mas que, exatamente por essa razão, é irrelevante para a ciência. Ora, o único objetivo que os filósofos têm em vista ao condenar os fenômenos é persuadir a si mesmos e aos outros de que o mundo real é muito diferente do mundo da aparência. Podemos todos compreender seu desejo de provar uma proposição tão desejável, mas não podemos congratulá-los por seu sucesso. É verdade que muitos deles não afirmam correspondentes objetivos para os fenômenos, e estes escapam à argumentação citada. Os que afirmam correspondentes são, via de regra, muito reticentes sobre o assunto, provavelmente porque sentem instintivamente que, se levado adiante, isso resultará num *rapprochement* muito grande entre o mundo fenomenal e o real. Se fossem explorar o tópico, dificilmente poderiam evitar a conclusão que estamos sugerindo. Sob esses aspectos, bem como em muitos outros, a noção de estrutura ou número de relação é importante.

Capítulo 7

Números racionais, reais e complexos

Vimos como definir números cardinais, e também números de relação, dos quais os comumente chamados números ordinais são uma espécie particular. Veremos que cada um desses tipos de número pode ser tanto infinito quanto finito. Mas nenhum deles é passível, sem alterações, das extensões mais conhecidas da idéia de número, a saber, as extensões para números negativos, fracionários, irracionais e complexos. No presente capítulo forneceremos brevemente definições lógicas dessas várias extensões.

Um dos erros que atrasaram a descoberta de definições corretas nessa região é a idéia comum de que cada extensão de número incluía as espécies anteriores como casos especiais. Pensava-se que, ao lidar com números inteiros positivos e negativos, os números inteiros positivos podiam ser identificados com os números inteiros originais sem sinal. Assim também, pensava-se que uma fração cujo denominador fosse 1 podia ser identificada com o número natural que é seu numerador. E supunha-se que os números irracionais, tal como a raiz quadrada de 2, encontravam seu lugar entre as frações racionais, como sendo maiores do que algumas delas e menores do que as outras, de modo que números racionais e irracionais podiam ser considerados conjuntamente uma classe, chamada “números reais”. E quando a idéia de número foi mais estendida de modo a incluir números “complexos”, isto é, números envolvendo a raiz quadrada de -1 , pensou-se que os números reais podiam ser considerados aqueles entre os números complexos em que a parte imaginária (isto é, a parte que era um múltiplo da raiz quadrada de -1) era zero. Todas essas suposições eram errôneas, e devem ser rejeitadas, como veremos, se quisermos chegar a definições corretas.

Começemos com números inteiros *positivos* e *negativos*. Um momento de consideração deixa óbvio que $+1$ e -1 devem ambos ser relações; de

fato, devem ser o inverso um do outro. A definição óbvia e suficiente é que $+1$ é a relação de $n + 1$ com n , e -1 é a relação de n com $n + 1$. Em geral, se m for um número indutivo, $+m$ será a relação de $n+m$ com n (para qualquer n), e $-m$ será a relação de n com $n + m$. De acordo com essa definição, $+m$ é uma relação um-um contanto que n seja um número cardinal (finito ou infinito) e m é um número cardinal indutivo. Mas $+m$ não pode, em nenhuma circunstância, ser identificado com m , que não é uma relação, mas uma classe de classes. De fato, $+m$ é exatamente tão distinto de m quanto $-m$.

Frações são mais interessantes que números inteiros positivos ou negativos. Precisamos de frações para muitas finalidades, porém talvez de maneira mais óbvia para mensuração. Meu amigo e colaborador dr. A.N. Whitehead desenvolveu uma teoria das frações especialmente adaptada para aplicação à mensuração, que está exposta em *Principia Mathematica*.¹ Mas se for necessário apenas definir objetos que tenham as propriedades puramente matemáticas requeridas, essa finalidade pode ser atingida por um método mais simples, que adotaremos aqui. Definiremos a fração m/n como sendo aquela relação existente entre dois números indutivos, x, y quando $xn = ym$. Essa definição nos permite provar que m/n é uma relação um-um, contanto que nem m nem n sejam zero. E, é claro, n/m é a relação inversa a m/n .

A partir da definição mencionada fica claro que a fração $m/1$ é a relação entre dois números inteiros x e y que consiste no fato de que $x = my$. Esta relação, como a relação $+m$, não pode de maneira alguma ser identificada com o número cardinal indutivo m , porque uma relação e uma classe de classes são objetos de tipos completamente diferentes.² Veremos que $0/n$ é sempre a mesma relação, não importa que número indutivo n seja; é, em suma, a relação de 0 com qualquer outro número inteiro indutivo cardinal. Podemos chamar isso o zero dos números racionais; ele não é, evidentemente, idêntico ao número cardinal 0. Inversamente, a relação $m/0$ é sempre a mesma, não importa que número indutivo m seja. Não há nenhum cardinal indutivo que corresponda a $m/0$. Podemos chamá-la “a infinidade dos racionais”. Trata-se de um caso do tipo de infinito tradicional na matemática, que é representado por “ ∞ ”. Esse totalmente diferente do verdadeiro infinito cantoriano, que consideraremos no próximo capítulo. A infinidade dos racionais não requer, para sua definição ou uso, quaisquer classes infinitas de números inteiros infinitos.

Não é, na verdade, uma noção muito importante, e poderíamos prescindir totalmente dela, se tivéssemos alguma razão para isso. O infinito cantoriano, por outro lado, é da maior e mais fundamental importância; a compreensão dele abre caminho para domínios inteiramente novos da matemática e da filosofia.

Convém observar que zero e infinidade são as únicas razões que não são um-um. Zero é um-muitos e infinidade é muitos-um.

Não há nenhuma dificuldade em definir *maior* e *menor* entre razões (ou frações). Dadas duas razões m/n e p/q , diremos que m/n é menor do que p/q se mq for menor do que pn . Não há nenhuma dificuldade em provar que a relação “menor do que”, assim definida, é serial, de modo que as razões formam uma série em ordem de magnitude. Nessa série, zero é o menor termo e infinidade é o maior. Se omitirmos zero e infinidade de nossa série, não haverá mais nenhuma razão menor ou maior; é óbvio que se m/n for uma razão diferente de zero e infinidade, $m/2n$ será menor e $2m/n$ será maior, embora nem um nem outro seja zero ou infinidade, de modo que m/n não será nem a menor nem a maior razão, e portanto (quando zero e infinidade forem omitidos) não haverá menor ou maior, uma vez que m/n foi escolhido arbitrariamente. De maneira semelhante, podemos provar que, por menos diferentes que duas frações possam ser, haverá sempre outras frações entre elas. Suponhamos que m/n e p/q são duas frações, das quais p/q é a maior. É fácil ver (ou provar) então que $(m + p) / (n + q)$ será maior do que m/n e menor do que p/q . Assim, a série de razões é tal que dois termos nunca são consecutivos, havendo sempre outros termos entre quaisquer dois. Como há outros termos entre esses outros, e assim por diante *ad infinitum*, é óbvio que há um número infinito de razões entre quaisquer duas, por menos diferentes que elas sejam.³ Uma série que possua a propriedade segundo a qual há sempre outros termos entre quaisquer dois, de modo que dois termos nunca sejam consecutivos, é chamada “compacta”. Portanto, as razões em ordem de magnitude formam uma série “compacta”. Essas séries têm importantes propriedades, e é importante observar que as razões fornecem um exemplo de série compacta gerada de maneira puramente lógica, sem nenhum apelo a espaço e tempo ou qualquer outro dado empírico.

Razões positivas e negativas podem ser definidas de maneira análoga àquela como definimos números inteiros positivos e negativos. Tendo primeiro definido a soma de duas razões m/n e p/q como $(mq + pn)/nq$,

definimos $+p/q$ como a relação de $m/n + p/q$ com m/n , em que m/n é uma razão; e $-p/q$ é, obviamente, o inverso de $+p/q$. Essa não é a única maneira possível de definir razões positivas e negativas, mas, para nossa finalidade, ela tem o mérito de ser uma adaptação óbvia daquela que adotamos no caso dos números inteiros.

Chegamos agora a uma extensão mais interessante da idéia de número, isto é, a extensão aos chamados números “reais”, que são o tipo que abarca os irracionais. No Capítulo 1 tivemos a oportunidade de mencionar “incomensuráveis” e sua descoberta por Pitágoras. Foi por meio deles, isto é, por meio da geometria, que se pensou pela primeira vez nos números irracionais. Um quadrado cujo lado tem um centímetro de comprimento terá uma diagonal cujo comprimento é a raiz quadrada de dois centímetros. Mas, como os antigos descobriram, não há nenhuma fração cujo quadrado seja 2. Essa proposição está provada no décimo livro de Euclides, que é um daqueles livros que estudantes supunham estar perdido, para a sua alegria, no tempo em que Euclides ainda era usado como livro-texto. A prova é extraordinariamente simples. Se possível, suponhamos que m/n é a raiz quadrada de 2, de modo que $m^2/n^2 = 2$, isto é, $m^2 = 2n^2$. Assim, m^2 é um número par, e portanto m deve ser um número par, porque o quadrado de um número ímpar é ímpar. Ora, se m é par, m^2 deve ser divisível por 4, pois se $m = 2p$, então $m^2 = 4p^2$. Assim teremos $4p^2 = 2n^2$, onde p é metade de m . Portanto, $2p^2 = n^2$, e por conseguinte n/p será também a raiz quadrada de 2. Mas podemos então repetir o raciocínio: se $n = 2q$, p/q será também a raiz quadrada de 2, e assim por diante, através de uma série interminável de números dos quais cada um é a metade de seu predecessor. Mas isso é impossível; se dividirmos um número por 2, e depois dividirmos a metade ao meio e assim por diante, atingiremos necessariamente um número ímpar após um número finito de passos. Podemos também formular o raciocínio de maneira ainda mais simples, supondo que o m/n com que começamos está em seus termos mais baixos; nesse caso, m e n não podem ambos ser pares; vimos, contudo, que, se $m^2/n^2 = 2$, devem ser. Não pode, haver portanto, nenhuma fração m/n cujo quadrado seja 2.

Por conseguinte, nenhuma fração expressará exatamente o comprimento da diagonal de um quadrado cujo lado tem um centímetro de comprimento. Isso parece um desafio lançado pela natureza para a aritmética. Por mais que o matemático se gabe (como o fez Pitágoras)

acerca do poder dos números, a natureza parece frustrá-lo exibindo comprimentos que número algum pode estimar em termos da unidade. Mas o problema não permaneceu nessa forma geométrica. Assim que a álgebra foi inventada, o mesmo problema surgiu com relação à solução de equações, embora aqui ele tenha assumido uma forma mais ampla, uma vez que envolvia também números complexos.

É claramente possível encontrar frações que cheguem cada vez mais perto de ter seu quadrado igual a 2. Podemos formar uma série ascendente de frações que tenham todas elas seus quadrados menores que do 2, mas diferindo de 2 e seus membros posteriores por menos do que um valor designado- α . Isto é, suponhamos que designo algum pequeno valor de antemão, digamos um bilionésimo, e será constatado que todos os termos de nossa série após um certo termo, digamos o décimo, têm quadrados que diferem de 2 por menos que esse valor. Se eu designasse um valor ainda menor, poderia ser necessário ir adiante ao longo da série, mas chegaríamos mais cedo ou mais tarde a um termo na série, digamos o vigésimo, após o qual todos os termos teriam quadrados diferindo de 2 por menos que esse valor ainda menor. Se tentássemos extrair a raiz quadrada de 2 pela regra aritmética usual, obteríamos um decimal interminável que, levado para tais ou tais lugares, preencheria exatamente as condições citadas. Podemos igualmente formar uma série descendente de frações cujos quadrados sejam todos maiores do que 2, mas maior por quantidades continuamente menores, à medida que chegamos aos termos posteriores da série, e diferindo, mais cedo ou mais tarde, por menos do que qualquer quantidade designada. Dessa maneira, é como se estivéssemos estendendo um cordão de isolamento em torno da raiz quadrada de 2, e talvez seja difícil acreditar que ela possa nos escapar permanentemente. Não é por esse método, porém, que chegaremos realmente à raiz quadrada de 2.

Se dividirmos *todas* as razões em duas classes, conformemente seus quadrados sejam menores do que 2 ou não, veremos que, entre aquelas cujos quadrados *não* são menores do que 2, todas têm seus quadrados maiores do que 2. Não há um máximo para as razões cujo quadrado é menor do que 2, e um mínimo para aquelas cujo quadrado é maior do que 2. Não há limite inferior, com exceção de zero, para as diferenças entre os números cujo quadrado é um pouco menor do que 2 e aqueles cujo quadrado é um pouco maior do que 2. Podemos, em suma, dividir *todas* as razões em duas classes tais que todos os termos em uma delas sejam

menores do que todos os termos na outra, não havendo máximo para uma classe e não havendo mínimo para a outra. Entre essas duas classes, onde $\sqrt{2}$ deveria estar, não há nada. Portanto, nosso cordão de isolamento, embora o tenhamos estreitado tanto quanto possível, foi estendido no lugar errado, e não prendeu $\sqrt{2}$.

O método descrito de dividir todos os termos de uma série em duas classes, uma das quais precede inteiramente a outra, foi posto em destaque por Dedekind,⁴ sendo por isso chamado “corte de Dedekind”. Com respeito ao que acontece no ponto de seção, há duas possibilidades: (1) pode haver um máximo para a seção inferior e um mínimo para a seção superior, (2) pode haver um máximo para uma e nenhum mínimo para a outra, (3) pode não haver um máximo para uma, mas um mínimo para a outra, (4) pode não haver um máximo para uma, nem um mínimo para a outra. Desses quatro casos, o primeiro é ilustrado por séries em que há termos consecutivos: na série dos números inteiros, por exemplo, uma seção inferior deve terminar com algum número n e a seção superior deve então começar com $n+1$. O segundo caso será ilustrado na série das razões se tomarmos como nossa seção inferior todas as razões até 1, incluindo 1, e em nossa seção superior todas as razões maiores do que 1. O terceiro caso é ilustrado se tomarmos por nossa seção inferior todas as razões menores do que 1 e por nossa seção superior todas as razões de 1 para cima (incluindo o próprio 1). O quarto caso, como vimos, é ilustrado se pusermos em nossa seção inferior todas as razões cujo quadrado for menor do que 2 e em nossa seção superior todas cujo quadrado for maior do que 2.

Podemos deixar de lado o primeiro de nossos quatro casos, visto que ele só surge em séries em que há termos consecutivos. No segundo de nossos quatro casos, dizemos que o máximo da seção inferior é o *limite inferior* da seção superior, ou de qualquer conjunto de termos escolhido na seção superior, de tal maneira que nenhum termo da seção superior esteja antes de todos eles. No terceiro de nossos quatro casos, dizemos que o mínimo da seção superior é o limite superior da seção inferior, ou de qualquer conjunto de termos escolhido na seção inferior, de tal maneira que nenhum termo da seção inferior esteja depois de todos eles. No quarto caso, dizemos que há uma “lacuna”: nem a seção superior nem a inferior têm um limite ou um último termo. Nesse caso, podemos também dizer que

temos uma “seção irracional”, visto que seções da série de razões têm “lacunas” quando correspondem a irracionais.

O que atrasou a verdadeira teoria dos irracionais foi uma crença errônea de que devia haver “limites” de séries de razões. A noção de “limite” é da maior importância, e antes de seguirmos adiante, é conveniente defini-la.

Diz-se que um termo x é um “limite superior” de uma classe α com respeito a uma relação P se (1) α não tiver nenhum máximo em P , (2) todos os membros de α que pertencem ao campo de P precedem x , (3) todos os membros do campo de P que precedem x precedem algum membro de α . (Por “precede” entendemos “tem a relação P com”.)

Isso pressupõe a seguinte definição de um “máximo”: diz-se que um termo x é um “máximo” de uma classe α com respeito a uma relação P se x for membro de α e do campo de P e não tiver a relação P com nenhum outro membro de α .

Essas definições não exigem que os termos a que são aplicadas sejam quantitativos. Por exemplo, dada uma série de momentos de tempo arranjados em anteriores e posteriores, seu “máximo” (se houver algum) será o último dos momentos; mas se eles estiverem arranjados em posteriores e anteriores, seu “mínimo” (se houver algum) será o primeiro dos momentos.

O “mínimo” de uma classe com respeito a P é seu máximo com respeito ao inverso de P ; e o “limite inferior” com respeito a P é o limite superior com respeito ao inverso de P .

As noções de limite e máximo não exigem essencialmente que a relação a respeito da qual são definidas seja serial, mas têm poucas aplicações importantes exceto a casos em que a relação é serial ou quase-serial. Uma noção muitas vezes importante é a de “limite superior ou máximo”, a que podemos chamar “fronteira superior”. Assim, a “fronteira superior” de um conjunto de termos escolhido numa série é seu último membro, se eles tiverem um, mas, se não, é o primeiro termo depois de todos eles, se houver tal termo. Se não houver nem um máximo nem um limite, não haverá fronteira superior. A “fronteira inferior” é um limite inferior ou mínimo.

Retornando aos quatro tipos de seção de Dedekind, vemos que no caso dos três primeiros tipos, cada seção tem uma fronteira (superior ou inferior, conforme o caso), ao passo que no quarto tipo nenhuma das duas tem fronteira. É claro também que, sempre que a seção inferior tiver uma

fronteira superior, a seção superior terá uma fronteira inferior. No segundo e terceiro casos, as duas fronteiras são idênticas; no primeiro, são termos consecutivos da série.

Uma série é chamada “dedekindiana” quando todas as seções têm uma fronteira, superior ou inferior conforme o caso. Vimos que a série das razões em ordem de magnitude não é dedekindiana.

Por força do hábito de se deixarem influenciar pela imaginação espacial, as pessoas supunham que as séries *deviam* ter limites em casos em que pareceria estranho que não tivessem. Assim, ao perceber que não havia nenhum limite racional para as razões cujo quadrado era menor do que 2, elas se permitiram “postular” um limite *irracional*, que deveria preencher a lacuna dedekindiana. Dedekind, na obra mencionada anteriormente, estabeleceu o axioma de que a lacuna deve sempre ser preenchida, isto é, que toda seção deve ter uma fronteira. É por essa razão que séries em que seu axioma é verificado são chamadas “dedekindianas”. Mas há um número infinito de séries para as quais ele não se verifica.

O método de “postular” o que queremos tem muitas vantagens; as mesmas vantagens do roubo sobre o trabalho honesto. Vamos deixá-lo para outros e seguir adiante com nossa labuta honesta.

É claro que um corte de Dedekind irracional “representa” de alguma maneira um irracional. Para fazer uso disso, que para começar não passa de uma vaga impressão, devemos encontrar alguma maneira de extrair dele uma definição precisa; e para fazê-lo, temos de livrar nossas mentes da noção de que um irracional deve ser o limite de um conjunto de razões. Assim como razões cujo denominador é 1 não são idênticas a números inteiros, assim também aqueles números racionais que podem ser maiores ou menores do que irracionais, ou podem ter irracionais como seus limites, não devem ser identificados com razões. Temos de definir um novo tipo de números, chamados “números reais”, dos quais alguns serão racionais e alguns irracionais. Os que são racionais “correspondem” a razões, do mesmo modo que a razão $n/1$ corresponde ao número inteiro n ; mas não são o mesmo que razões. Para decidir o que eles devem ser, observemos que um irracional é representado por um corte irracional, e um corte é representado por sua seção inferior. Limitemo-nos a cortes em que a seção inferior não tenha nenhum máximo; nesse caso chamaremos a seção inferior um “segmento”. Portanto aqueles segmentos que correspondem a razões são os que consistem em todas as razões menos a

razão a que eles correspondem, a qual é seu limite; enquanto aqueles que representam irracionais são os que não têm nenhuma fronteira. Os segmentos, tanto os que têm fronteiras quanto os que não têm, são tais que, de quaisquer dois pertencentes a uma série, um deve ser parte do outro; em consequência eles podem todos ser arranjados numa série pela relação de todo e parte. Uma série em que haja lacunas de Dedekind, isto é, em que haja segmentos que não têm fronteira, dará origem a um número de segmentos maior do que o número de seus termos, uma vez que cada termo definirá um segmento tendo aquele termo por fronteira, e assim os segmentos sem fronteira serão extras.

Estamos agora em condições de definir número real e número irracional. Um “número real” é um segmento da série de razões em ordem de magnitude. Um “número irracional” é um segmento da série de razões que não tem nenhuma fronteira. Um “número real racional” é um segmento da série de razões que tem uma fronteira.

Um número real racional consiste, portanto, em todas as razões menores do que determinada razão, e é o número real racional correspondente a essa razão. O número real 1, por exemplo, é a classe das frações próprias.

Nos casos em que supusemos naturalmente que um irracional devia ser o limite de um conjunto de razões, a verdade é que ele é o limite do conjunto correspondente de números reais racionais na série de segmentos ordenados por todo e parte. Por exemplo, $\sqrt{2}$ é o limite superior daqueles segmentos da série de razões que correspondem a razões cujo quadrado é menor que 3. Mais simplesmente ainda, $\sqrt{2}$ é o segmento que *consiste* em todas aquelas razões cujo quadrado é menor que 2.

É fácil provar que a série de segmentos de qualquer série é dedekindiana. Pois, dado qualquer conjunto de segmentos, sua fronteira será sua soma lógica, isto é, a classe de todos aqueles termos que pertencem a pelo menos um segmento do conjunto.⁵

A definição de número real que acabamos de citar é um exemplo de “construção” em contraposição a “postulação”, de que tivemos outro exemplo na definição de números cardinais. A grande vantagem desse método é não requerer nenhuma nova suposição, permitindo-nos proceder dedutivamente a partir do aparato original da lógica.

Não há nenhuma dificuldade em definir adição e multiplicação para números reais como definidos previamente. Dados dois números reais, μ e ν , cada um sendo uma classe de razões, tome qualquer membro de μ e

qualquer membro de v e some-os segundo a regra para a adição de razões. Forme a classe de todas aquelas somas obteníveis variando-se os números selecionados de μ e v . Isso dá uma nova classe de razões, e é fácil provar que essa nova classe é um segmento da série de razões. Nós a definimos como a soma de μ e v . Podemos expressar a definição mais brevemente da seguinte maneira: *a soma aritmética de dois números reais é a classe das somas aritméticas de um membro de uma e um membro da outra escolhidos de todas as maneiras possíveis.*

Podemos definir o produto aritmético de dois números reais exatamente da mesma maneira, multiplicando um membro de uma por um membro da outra de todas as maneiras possíveis. A classe de razões assim gerada é definida como o produto dos dois números reais. (Em todas essas definições, a série de razões deve ser definida de modo a excluir 0 e a infinidade.)

Não há dificuldade em estender nossas definições a números reais positivos e negativos e à sua adição e multiplicação. Falta dar a definição de números complexos.

Os números complexos, embora passíveis de uma interpretação geométrica, não são exigidos pela geometria da mesma maneira imperativa como o são os números irracionais. Um número “complexo” significa um número que envolva a raiz quadrada de um número negativo, quer seja integral, fracionário ou real. Como o quadrado de um número negativo é positivo, um número cujo quadrado deve ser negativo tem de ser um novo tipo de número. Usando a letra i para a raiz quadrada de -1 , qualquer número que envolva a raiz quadrada de um número negativo pode ser expressa na forma $x + yi$, onde x é real. A parte yi é chamada a parte “imaginária” desse número, x sendo a parte “real”. (A razão para a expressão “números reais” é que eles são contrapostos aos “imaginários”.) Números complexos foram habitualmente usados por matemáticos por um longo tempo apesar da ausência de qualquer definição precisa. Supunha-se simplesmente que deviam obedecer às regras matemáticas usuais, e com base nesses pressupostos considerou-se que seu uso era vantajoso. Eram requeridos menos para geometria do que para álgebra e análise. Desejamos, por exemplo, ser capazes de dizer que toda equação quadrática tem duas raízes, e toda equação cúbica tem três, e assim por diante. Mas se ficarmos restritos aos números reais, uma equação como $x^2 + 1 = 0$ não tem raiz nenhuma, e uma equação como $x^3 - 1 = 0$ tem apenas uma. Toda

generalização de número apresentou-se primeiro como necessária para algum problema simples: números negativos foram necessários para que a subtração pudesse ser sempre possível, visto que de outro modo $a - b$ seria sem sentido se a fosse menor do que b ; frações foram necessárias para que a divisão pudesse ser sempre possível; e números complexos são necessários para que extração de raízes e solução de equações possam ser sempre possíveis. Mas extensões de número não são *criadas* pela mera necessidade deles: são criadas pela definição, e é para a definição de números complexos que voltaremos agora nossa atenção.

Um número complexo pode ser considerado e definido como simplesmente um par ordenado de números reais. Aqui, como em outros lugares, muitas definições são possíveis. É necessário apenas que a definição adotada conduza a certas propriedades. No caso de números complexos, se eles forem definidos como pares ordenados de números reais, teremos assegurado de imediato algumas das propriedades requeridas, a saber, que dois números reais são necessários para determinar um número complexo, e que entre esses podemos distinguir um primeiro e um segundo, e que dois números complexos só serão idênticos quando o primeiro número real envolvido em um for igual ao primeiro envolvido no outro, e o segundo ao segundo. O que é necessário além disso pode ser assegurado mediante a definição das regras de adição e multiplicação. Devemos ter

$$\begin{aligned}(x + yi) + (x' + y'i) &= (x + x') + (y + y')i \\ (x + yi)(x' + y'i) &= (xx' - yy') + (xy' + x'y)i.\end{aligned}$$

Definiremos assim que, dados dois pares ordenados de números reais, (x, y) e (x', y') , sua soma deve ser o par $(x + x', y + y')$, e seu produto deve ser o par $(xx' - yy', xy' + x'y)$. Com essas definições garantiremos que nossos pares ordenados terão as propriedades que desejamos. Por exemplo, tomemos o produto dos dois pares $(0, y)$ e $(0, y')$. Esse será, pela regra citada, o par $(-yy', 0)$. Assim, o quadrado do par $(0, 1)$ será o par $(-1, 0)$. Ora, os pares em que o segundo termo é 0 são aqueles que, segundo a nomenclatura usual, têm sua parte imaginária zero; na notação $x + yi$, eles são $x + 0i$, que é natural escrever simplesmente x . Assim como é natural (mas errôneo) identificar razões cujo denominador é a unidade com números inteiros, assim também é natural (mas errôneo) identificar

números complexos cuja parte imaginária é zero com números reais. Embora isso seja um erro na teoria, é uma conveniência na prática; “ $x + 0i$ ” pode ser substituído simplesmente por “ x ” e “ $0 + yi$ ” por “ yi ”, contanto que lembremos que o “ x ” não é realmente um número real, mas um caso especial de um número complexo. E quando y é 1, “ yi ” pode, é claro, ser substituído por “ i ”. Assim o par $(0, 1)$ é representado por i , e o par $(-1, 0)$ é representado por -1 . Ora, nossas regras de multiplicação tornam o quadrado de $(0, 1)$ igual a $(-1, 0)$, isto é, o quadrado de i é -1 . Era isso que desejávamos assegurar. Nossas definições servirão, portanto, a todas as finalidades necessárias.

É fácil dar uma interpretação geométrica de números complexos na geometria do plano. Esse tema foi agradavelmente exposto por W.K. Clifford em seu *Common Sense of the Exact Sciences*, um livro de grande mérito, mas escrito antes que a importância de definições puramente lógicas fosse compreendida.

Números complexos de uma ordem mais elevada, embora muito menos úteis e importantes que os que acabamos de definir, têm certos usos não de todo desprovidos de importância na geometria, como pode ser visto, por exemplo, em *Universal Algebra* do dr. Whitehead. A definição de números complexos da ordem n é obtida mediante uma extensão óbvia da definição que demos. Definimos um número complexo da ordem n como uma relação um-muitos cujo domínio consiste em certos números reais e cujo domínio inverso consiste nos números inteiros de 1 a n .^{*} Isso é o que seria comumente indicado pela notação $(x_1, x_2, x_3, \dots, x_n)$, onde os sufixos denotam correlação com os números inteiros usados como sufixos, e a correlação é um-muitos, não necessariamente um-um, porque x_r e x_s podem ser iguais quando r e s não são iguais. A definição mencionada, com uma regra de multiplicação adequada, servirá a todas as finalidades para as quais números complexos de ordens mais elevadas são necessários.

Completamos nossa revisão das extensões de número que não envolvem infinidade. A aplicação de número a coleções infinitas deve ser nosso próximo tópico.

^{*} Cf. *Principles of Mathematics*, § 360, p.379.

Capítulo 8

Números cardinais infinitos

A definição de números cardinais que demos no Capítulo 2 foi aplicada no Capítulo 3 a números finitos, isto é, aos números naturais ordinários. Demos a esses o nome “números indutivos”, porque verificamos que devem ser definidos como números que obedecem à indução matemática, começando por 0. Mas ainda não consideramos coleções que não têm um número indutivo de termos, nem indagamos se é possível dizer que tais coleções têm de fato um número. Esse é um problema antigo, que foi resolvido em nosso próprio tempo, principalmente por Georg Cantor. No presente capítulo tentaremos explicar a teoria dos números cardinais transfinitos ou infinitos tal como ela resulta de uma combinação das descobertas de Cantor com as de Frege no campo da teoria lógica dos números.

Não se pode dizer que é *certo* que existem coleções infinitas no mundo. A suposição de que existem é chamada “axioma da infinidade”. Embora haja aparentemente várias maneiras pelas quais poderíamos esperar provar esse axioma, há razões para se temer que sejam todas falaciosas, e que não haja razão lógica conclusiva para se acreditar que ele é verdadeiro. Ao mesmo tempo, certamente não existe razão lógica alguma *contra* coleções infinitas, e por isso estamos logicamente justificados ao investigar a hipótese de que tais coleções existam. A forma prática dessa hipótese, para nossas finalidades presentes, é a suposição de que, se n for qualquer número indutivo, n não será igual a $n + 1$. Várias sutilezas surgem quando identificamos essa forma de nossa suposição com a forma que afirma a existência de coleções infinitas; mas nós as deixaremos de lado até passarmos, num capítulo posterior, a considerar o axioma da infinidade por si mesmo. Por enquanto, vamos simplesmente supor que, se n for um número indutivo, n não será igual a $n + 1$. Isso está envolvido na suposição feita por Peano de que dois números indutivos diferentes jamais têm o

mesmo sucessor; pois, se $n = n + 1$, então $n + 1$ e n têm o mesmo sucessor, a saber, n . Não estamos, portanto, supondo nada que não estivesse envolvido nas proposições primitivas de Peano.

Consideremos agora a coleção dos próprios números indutivos. Essa é uma classe perfeitamente bem definida. Em primeiro lugar, um número cardinal é um conjunto de classes, todas similares umas às outras e não similares a nada exceto umas às outras. Definimos então como “números indutivos” aqueles entre os cardinais que pertencem à posteridade de 0 com respeito à relação de n com $n + 1$, isto é, aqueles que possuem todas as propriedades de 0 e seus sucessores, entendendo por “sucessor” de n o número $n + 1$. Assim a classe dos “números indutivos” é perfeitamente definida. Por nossa definição geral de números cardinais, o número de termos na classe dos números indutivos deve ser “todas as classes similares à classe dos números indutivos”, isto é, esse conjunto de classes é o número dos números indutivos segundo nossas definições.

É fácil ver agora que esse número não é um dos números indutivos. Se n for um número indutivo, o número de números de 0 a n (ambos incluídos) será $n + 1$; portanto, o número total de números indutivos será maior do que n , não importando qual dos números indutivos n possa ser. Se arranjarmos os números indutivos numa série em ordem de magnitude, esta série não terá nenhum último termo; mas se n for um número indutivo, toda série cujo campo tiver n termos terá um último termo, como é fácil provar. Tais diferenças poderiam ser multiplicadas *ad libitum*. Assim, o número de números indutivos é um novo número, diferente de todos eles, não possuindo todas as propriedades indutivas. Pode acontecer que 0 tenha determinada propriedade, e que se n a tiver, $n + 1$ também a terá; esse novo número, contudo, não a terá. As dificuldades que atrasaram por tanto tempo a teoria dos números infinitos deveram-se em grande parte ao fato de se considerar que necessariamente todos os números possuem pelo menos algumas propriedades indutivas; de fato, pensava-se que elas não poderiam ser negadas sem contradição. O primeiro passo para compreendermos números infinitos consiste em perceber esse erro.

A diferença mais digna de nota e espantosa entre um número indutivo e esse novo número é que este permanece inalterado pela adição ou subtração de 1; permanecerá igualmente inalterado se o dobrarmos, ou o dividirmos por 2 ou o submetermos a qualquer das outras operações que nos parecem tornar um número necessariamente maior ou menor. O fato

de não ser alterado pela adição de 1 é usado por Cantor para a definição do que ele chama números cardinais “transfinitos”; mas, por várias razões, algumas das quais emergirão à medida que prosseguirmos, é melhor definir um número cardinal infinito como um número que não possui todas as propriedades indutivas, isto é, simplesmente um número que não é um número indutivo. Apesar disso, a propriedade de não ser alterado pela adição de 1 é muito importante e devemos nos entender um pouco sobre ela.

Dizer que uma classe tem um número que não é alterado pela adição de 1 é o mesmo que dizer que, se tomarmos um termo x que não pertence à classe, podemos encontrar uma relação um-um cujo domínio é a classe e cujo domínio inverso é obtido somando-se x à classe. Pois, nesse caso, a classe é similar à soma de si mesma e do termo x , isto é, a uma classe que tem um termo extra; de modo que ela tem o mesmo número que uma classe com um termo extra, de tal modo que, se n for esse número, $n = n + 1$. Nesse caso, teremos também $n = n - 1$, isto é, haverá relações um-um cujo domínio consiste em toda a classe e cujo domínio inverso consiste em apenas um termo, excluída toda a classe. Pode ser mostrado que os casos em que isso acontece são aqueles mesmos casos aparentemente mais gerais em que *alguma* parte (excluído o todo) pode ser posta em relação um-um com o todo. Quando isso pode ser feito, pode-se dizer que o correlator pelo qual é feito “reflete” toda a classe numa parte dela mesma; por essa razão, essas classes podem ser chamadas “reflexivas”. Assim: uma classe “reflexiva” é uma classe similar a uma parte própria de si mesma. (Uma “parte própria” é uma parte que não é o todo.)

Um número cardinal “reflexivo” é o número cardinal de uma classe reflexiva.

Devemos agora considerar essa propriedade da reflexividade. Um dos casos mais impressionantes de “reflexão” é o exemplo do mapa, de Royce. Ele imagina que se decidiu traçar um mapa da Inglaterra numa parte da superfície da Inglaterra. Um mapa, para ser preciso, deve ter uma correspondência um-um perfeita com seu original; assim, nosso mapa, que é parte, é uma relação um-um com o todo, e deve conter o mesmo número de pontos que o todo, o qual deve portanto ser um número reflexivo. Royce está interessado no fato de que o mapa, para ser correto, deve conter um mapa do mapa, o qual por sua vez deve conter um mapa do mapa do mapa, e assim *ad infinitum*. A idéia é interessante, mas não precisa nos

ocupar neste momento. De fato, faremos bem em passar de ilustrações pitorescas a outras mais completamente definidas, e para esse propósito não há nada melhor do que considerar a própria série dos números.

A relação de n com $n + 1$, limitada a números indutivos, é um-um, tem a totalidade dos números indutivos por seu domínio, e todos, exceto 0, por seu domínio inverso. Assim, a totalidade da classe dos números indutivos é similar ao que a mesma classe se torna quando omitimos 0. Conseqüentemente, ela é uma classe “reflexiva” segundo a definição, e o número de seus termos é um número “reflexivo”. Novamente, a relação de n com $2n$, limitada a números indutivos, é um-um, tem a totalidade dos números indutivos por domínio e os números indutivos pares apenas por seu domínio inverso. Portanto, o número total de números indutivos é o mesmo que o número de números indutivos pares. Essa propriedade foi usada por Leibniz (e muitos outros) como prova de que os números infinitos são impossíveis; considerava-se contraditório que “a parte devesse ser igual ao todo”. Mas essa é uma daquelas expressões cuja plausibilidade depende de uma imprecisão não percebida: a palavra “igual” tem muitos sentidos, mas se for tomada como “similar”, não haverá contradição, uma vez que uma coleção infinita pode perfeitamente ter partes similares a si mesma. Aqueles que consideraram isso impossível, atribuíram aos números em geral, inconscientemente, propriedades que só podem ser provadas por indução matemática, e que somos levados, apenas por sua familiaridade, a considerar, erroneamente, verdadeiras além da região do finito.

Sempre que podemos “refletir” uma classe numa parte de si própria, a mesma relação refletirá necessariamente numa parte menor, e assim por diante *ad infinitum*. Por exemplo, podemos refletir, como acabamos de ver, todos os números indutivos nos números pares; podemos, pela mesma relação (a de n com $2n$), refletir os números pares nos múltiplos de 4, estes nos múltiplos de 8, e assim por diante. Esse é um análogo abstrato do problema do mapa de Royce. Os números pares são um “mapa” de todos os números indutivos; os múltiplos de 4 são um mapa do mapa; os múltiplos de 8 são um mapa do mapa do mapa, e assim por diante. Se tivéssemos aplicado o mesmo processo à relação de n com $n + 1$, nosso “mapa” teria consistido em todos os números indutivos exceto 0; o mapa do mapa teria consistido em todos a partir de 2, o mapa do mapa do mapa em todos a partir de 3, e assim por diante. O principal uso dessa ilustração

é nos familiarizar com a idéia de classes reflexivas, de tal modo que proposições aritméticas aparentemente paradoxais possam ser prontamente traduzidas na linguagem de reflexões e classes, em que a aparência de paradoxo é muito menor.

Será útil dar uma definição de número que seja a dos cardinais indutivos. Para essa finalidade, definiremos primeiro o tipo de série exemplificado pelos cardinais indutivos em ordem de magnitude. O tipo de série chamado “progressão” já foi considerado no Capítulo 1. É uma série que pode ser gerada por uma relação de consecutividade: todos os números na série devem ter um sucessor, mas deve haver apenas um que não tenha predecessor, e todos os membros da série devem estar em posteridade a esse termo com respeito à relação “predecessor imediato”. Essas características podem ser resumidas na seguinte definição:¹ uma “progressão” é uma relação um-um tal que há apenas um termo pertencente ao domínio, mas não ao domínio inverso, e o domínio é idêntico à posteridade desse único termo.

É fácil ver que uma progressão, assim definida, satisfaz os cinco axiomas de Peano. O termo pertencente ao domínio, mas não ao domínio inverso, será o que ele chama “0”; o termo com o qual um termo tem a relação um-um será o “sucessor” do termo; e o domínio da relação um-um será o que ele chama “número”. Tomando seus cinco axiomas sucessivamente, temos as seguintes traduções:

(1) “0 é um número” torna-se: “O membro do domínio que não é membro do domínio inverso é membro do domínio.” Isso é equivalente à existência de tal membro, o qual é dado em nossa definição. Chamaremos a esse membro “o primeiro termo”.

(2) “O sucessor de todo número é um número” torna-se: “O termo com que um dado membro do domínio tem a relação em questão é novamente membro do domínio.” Isso é provado da seguinte maneira: pela definição, todo membro do domínio é membro da posteridade do primeiro termo; portanto, o sucessor de um membro do domínio deve ser membro da posteridade do primeiro termo (porque a posteridade de um termo sempre contém os sucessores dele próprio, pela definição geral de posteridade), e portanto membro do domínio, porque pela definição a posteridade do primeiro termo é o mesmo que o domínio.

(3) “Dois números diferentes nunca têm o mesmo sucessor.” Isso quer dizer apenas que a relação é um-muitos, o que ela é por definição (sendo

um-um).

(4) “0 não é o sucessor de nenhum número” torna-se: “O primeiro termo não é membro do domínio inverso”, o que é novamente um resultado imediato da definição.

(5) Isso é indução matemática e torna-se “todo membro do domínio pertence à posteridade do primeiro termo”, o que era parte de nossa definição.

Dessa maneira, as progressões, tais como as definimos, têm as cinco propriedades formais das quais Peano deduz a aritmética. É fácil mostrar que duas progressões são “similares” no sentido definido para similaridade de relações no Capítulo 6. Podemos, é claro, derivar uma relação serial da relação um-um pela qual definimos uma progressão: o método usado é aquele explicado no Capítulo 4, e a relação é aquela de um termo com um membro de sua posteridade própria com respeito à relação original um-um.

Duas relações assimétricas transitivas que geram progressões são similares, pelas mesmas razões pelas quais as relações um-um correspondentes são similares. A classe de todos esses geradores transitivos de progressões é um “número serial” no sentido do Capítulo 6; é de fato o menor dos números seriais infinitos, o número a que Cantor deu o nome ω , pelo qual o tornou famoso.

Mas o que nos interessa, por enquanto, são os números *cardinais*. Uma vez que duas progressões são relações similares, segue-se que seus domínios (ou seus campos, que são o mesmo que seus domínios) são classes similares. Os domínios das progressões formam um número cardinal, uma vez que se pode mostrar facilmente que toda classe similar ao domínio é ela mesma o domínio de uma progressão. Esse número cardinal é o menor dos números cardinais infinitos; é aquele a que Cantor chamou apropriadamente do hebraico álefe com o sufixo 0, para distingui-lo de cardinais infinitos maiores, que têm outros sufixos. Assim o nome do menor dos cardinais infinito é n_0 .

Dizer que uma classe tem ? termos é o mesmo que dizer que é membro de n_0 , e isso é o mesmo que dizer que os membros da classe podem ser arranjados numa progressão. É óbvio que toda progressão continua sendo uma progressão se omitimos dela um número finito de termos, ou um termo a cada dois, ou todos exceto cada décimo ou cada centésimo. Esses métodos de reduzir uma progressão não a fazem deixar de ser uma

progressão, e portanto não diminuem o número de seus termos, que continua sendo n_0 . De fato, qualquer seleção a partir de uma progressão é uma progressão se não tiver último termo, por mais dispersa que seja sua distribuição. Tomemos (digamos) números indutivos da forma n^n ou n^{nn} . Esses números vão ficando muito rarefeitos nas partes mais elevadas da série de números, mas apesar disso são exatamente tão numerosos quanto a totalidade dos números indutivos, a saber n_0 .

Inversamente, podemos acrescentar termos aos números indutivos sem aumentar seu número. Tomemos, por exemplo, razões. Poderíamos ter a tendência a pensar que deve haver muito mais razões do que números inteiros, visto que razões cujo denominador é 1 correspondem aos números inteiros e parecem ser apenas uma proporção infinitesimal das razões. Na realidade, porém, o número de razões (ou frações) é exatamente igual ao de números indutivos, a saber, n_0 . É fácil ver isso arranjando razões numa série de acordo com o seguinte plano: se a soma do numerador e do denominador em uma for menor do que na outra, põe-se uma antes da outra; se a soma for igual nas duas, põe-se primeiro a que tem menor numerador. Isso nos dá a série

$$1, 1/2, 2, 1/3, 3, 1/4, 2/3, 3/2, 4, 1/5, \dots$$

Essa série é uma progressão, e todas as razões ocorrem nela mais cedo ou mais tarde. Assim podemos arranjar as razões em progressão e portanto seu número é n_0 .

Não é verdade, no entanto, que *todas* as coleções infinitas têm n_0 termos. O número dos números reais, por exemplo, é maior do que n_0 ; ele é, de fato, 2^{n_0} , e não é difícil provar que 2^n é maior do que n mesmo quando n é infinito. A maneira mais fácil de provar isso é provar, em primeiro lugar, que se uma classe tiver n membros, ela contém 2^n subclasses — em outras palavras, que há 2^n maneiras de selecionar alguns de seus membros (incluindo os casos extremos em que selecionamos todos ou nenhum); e, em segundo lugar, do que o número de subclasses contido em uma classe é sempre maior que o número dos membros da classe. Dessas duas proposições, a primeira é conhecida no caso dos números finitos, e não é difícil estendê-la aos números infinitos. A prova da segunda é tão simples e tão instrutiva que vamos dá-la.

Em primeiro lugar, é claro que o número de subclasses de uma dada classe (digamos α) é pelo menos tão grande quanto o número de membros, visto que cada membro constitui uma subclasse; temos, portanto, uma correlação de todos os membros com algumas das subclasses. Segue-se portanto que, se o número de subclasses não for *igual* ao número de membros, deve ser *maior*. Ora, é fácil provar que o número não é igual, mostrando que, dada qualquer relação um-um cujo domínio sejam os membros e cujo domínio inverso esteja contido entre o conjunto de subclasses, deve haver pelo menos uma subclasse não pertencente ao domínio inverso. A prova é a seguinte:² quando uma correlação R um-um é estabelecida entre todos os membros de α e algumas das subclasses, pode acontecer que um dado membro x seja correlacionado a uma subclasse da qual é membro; ou ainda, pode acontecer que x seja correlacionado com uma subclasse da qual não é membro. Formemos toda a classe, digamos β , daqueles membros x correlacionados com subclasses de que não são membros. Essa é uma subclasse de α , e ela não é correlacionada com nenhum membro de α . Pois, tomando primeiro os membros de β , cada um deles (pela definição de β) é correlacionado com alguma subclasse da qual não é membro, e portanto não é correlacionado com β . Tomando em seguida os termos que não são membros de β , cada um deles (pela definição de β) é correlacionado com alguma subclasse da qual é membro, e portanto novamente não é correlacionado com β . Assim, nenhum membro de α é correlacionado com β . Como R era *qualquer* correlação um-um de todos os membros com algumas subclasses, segue-se que não há correlação de todos os membros com *todas* as subclasses. Não importa para essa prova que β não tenha nenhum membro: tudo o que ocorre nesse caso é que a subclasse que vemos ser omitida é a classe nula. Portanto, em qualquer caso, o número de subclasses não é igual ao número de membros, e por conseguinte, pelo que foi dito antes, é maior. Combinando isso com a proposição segundo a qual se n for o número de membros 2^n será o número de subclasses, temos o teorema segundo o qual 2^n será sempre maior do que n , mesmo quando n for infinito.

Segue-se dessa proposição que não há nenhum máximo para os números cardinais infinitos. Por maior que um número infinito n possa ser, 2^n será ainda maior. A aritmética dos números infinitos é um tanto surpreendente até que nos acostumemos a ela. Temos, por exemplo,

$$\begin{aligned}n_0 + 1 &= n_0 \\n_0 + n &= n_0, \text{ onde } n \text{ é qualquer número indutivo,} \\n_0^2 &= n_0.\end{aligned}$$

(Isso se segue do caso das razões, pois, visto que uma razão é determinada por um par de números indutivos, é fácil ver que o número de razões é o quadrado do número de números indutivos, isto é, é n_0^2 ; mas vimos que é também n_0 .)

$n_0^n = n_0$ onde n é qualquer número indutivo. (Isso se segue de $n_0^2 = n_0$ por indução; pois se $n_0^n = ?_0$, então $n_0^{n+1} = n_0^2 = n_0$.

Mas $2^n_0 > n_0$.

De fato, como veremos mais tarde 2^n_0 é um número muito importante, a saber, o número de termos numa série que tem “continuidade” no sentido em que essa palavra é usada por Cantor. Supondo que espaço e tempo são contínuos nesse sentido (como geralmente fazemos em geometria analítica e cinemática), esse será o número de pontos no espaço ou de instantes no tempo; será também o número de pontos em qualquer porção finita do espaço, seja ela linha, área ou volume. Depois de n_0 , 2^n_0 é o mais importante e interessante dos números cardinais.

Embora a adição e a multiplicação sejam sempre possíveis com cardinais infinitos, a subtração e a divisão deixam de dar resultados definidos, não podendo por isso ser empregadas como são em aritmética elementar. Tomemos a subtração para começar: enquanto o número subtraído for finito, tudo vai bem; se o outro número for reflexivo, ele permanece inalterado. Assim $n_0 - n = n_0$, se n for finito; até aí, a subtração dá um resultado perfeitamente definido. Mas é diferente quando subtraímos n_0 dele mesmo; podemos então obter qualquer resultado, de 0 até n_0 . Isso é igualmente visto por exemplos. Dos números indutivos, retire as seguintes coleções de termos $?_0$:

- (1) Todos os números indutivos — resto zero.
- (2) Todos os números indutivos de n em diante — resto os números de 0 a $n - 1$, numerando n termos ao todo.
- (3) Todos os números ímpares — resto, todos os números pares, numerando n_0 .

Todas essas são diferentes maneiras de subtrair n_0 de n_0 e todas dão resultados diferentes.

No tocante à divisão, resultados muitos similares seguem-se do fato de que n_0 fica inalterado quando multiplicado por 2 ou 3 ou qualquer número finito n ou por n_0 . Segue-se que n_0 dividido por n_0 pode ter qualquer valor de 1 a n_0 .

Da ambigüidade da subtração e da divisão resulta que números negativos e razões não podem ser estendidos a números infinitos. Adição, multiplicação e exponenciação procedem muito satisfatoriamente, mas as operações inversas — subtração, divisão e extração de raízes — são ambíguas, e as noções que dependem delas não se sustentam quando números infinitos estão envolvidos.

A característica pela qual definimos finitude foi indução matemática, isto é, definimos um número como finito quando ele obedece à indução matemática começando de 0, e uma classe como finita quando seu número é finito. Essa definição produz o tipo de resultado que uma definição deveria produzir, a saber, que os números finitos são aqueles que ocorrem nas séries de números comuns 0, 1, 2, 3, . . . Nesse capítulo, porém, os números infinitos que discutimos não foram meramente não-indutivos: foram também *reflexivos*. Cantor usou a reflexividade como a *definição* do infinito, e acredita que ela é equivalente à não-indutividade, isto é, acredita que toda classe e todo cardinal são ou indutivos ou reflexivos. Isso pode ser verdadeiro, e pode muito possivelmente ser passível de prova; mas as provas até agora oferecidas por Cantor e outros (incluindo o presente autor em dias passados) são falaciosas, por razões que serão explicadas quando chegarmos a considerar o “axioma multiplicativo”. No presente, não se sabe se há classes e cardinais que não sejam nem reflexivos nem indutivos. Se n fosse um cardinal desse tipo, não teríamos $n = n + 1$, mas n não seria um dos “números naturais” e careceria de algumas das propriedades indutivas. Todas as classes e cardinais infinitos *conhecidos* são reflexivos; mas, por enquanto, convém manter a mente aberta para a possibilidade de haver casos, até agora desconhecidos, de classes e cardinais que não sejam nem reflexivos nem indutivos. Enquanto isso, adotaremos as seguintes definições: Uma classe ou cardinal *finitos* são uma classe ou cardinal *indutivos*. Uma classe ou cardinal *infinitos* são uma classe ou cardinal *não indutivos*. Todas as classes e cardinais *reflexivos* são infinitos; mas não se sabe presentemente se todas as classes e

cardinais infinitos são reflexivos. Retornaremos a esse assunto no Capítulo 12.

Capítulo 9

Séries infinitas e ordinais

Uma “série infinita” pode ser definida como uma série cujo campo é uma classe infinita. Já tratamos de um tipo de série infinita, a saber, progressões. Neste capítulo vamos considerar o tema de maneira mais geral.

A característica mais notável de uma série infinita é que se pode alterar seu número serial meramente rearranjando seus termos. Sob esse aspecto, há certa contraposição entre números cardinais e seriais. É possível manter inalterado o número cardinal de uma classe reflexiva mesmo lhe acrescentando termos; por outro lado, é possível mudar o número serial de uma série sem lhe acrescentar nem lhe retirar nenhum termo, por mero rearranjo. Ao mesmo tempo, no caso de qualquer série infinita, é possível também, como com os cardinais, acrescentar termos sem alterar o número serial: tudo depende do modo como são acrescentados.

Para esclarecer o assunto, é melhor começar com exemplos. Consideremos primeiro vários tipos diferentes de série que podem ser formadas a partir dos números indutivos arranjados segundo vários planos. Começemos com a série

$$1, 2, 3, 4, \dots n, \dots,$$

que, como já vimos, representa o menor dos números seriais infinitos, o tipo que Cantor chama φ . Para refinar essa série, continuemos a efetuar repetidamente a operação de remover para o fim o primeiro número par que ocorrer. Obtemos assim em sucessão as várias séries:

$$\begin{aligned} &1, 3, 4, 5, \dots n, \dots 2, \\ &1, 3, 5, 6, \dots n+1, \dots 2, 4, \\ &1, 3, 5, 7, \dots n+2, \dots 2, 4, 6, \end{aligned}$$

e assim por diante. Se imaginarmos esse processo efetuado tantas vezes quanto possível, chegamos finalmente à série

$$1, 3, 5, 7, \dots 2n + 1, \dots 2, 4, 6, 8, \dots 2n \dots,$$

em que temos primeiro todos os números ímpares e depois todos os números pares.

Os números seriais dessas várias séries são $\varphi + 1, \varphi + 2, \varphi + 3, \dots 2\varphi$. Cada um desses números é “maior” do que qualquer de seus predecessores, no seguinte sentido. Diz-se que um número é “maior” do que outro se qualquer série que tenha o primeiro número contiver uma parte que tenha o segundo número, mas nenhuma série que tenha o segundo contiver uma parte que tenha o primeiro número.

Se compararmos as duas séries

$$\begin{array}{c} 1, 2, 3, 4, \dots n, \dots \\ 1, 3, 4, 5, \dots n + 1, \dots 2 \end{array}$$

veremos que a primeira é similar à parte da segunda que omite o último termo, a saber, o número 2, mas a segunda não é similar a nenhuma parte da primeira. (Isso é óbvio, mas pode ser facilmente demonstrado.) Assim, a segunda série tem um número serial maior do que a primeira, segundo a definição — isto é, $\varphi + 1$ é maior do que φ . Mas se acrescentarmos um termo no início de uma progressão em vez de no fim, ainda teremos uma progressão. Assim $1 + \varphi = \varphi$. Portanto, $1 + \varphi$ não é igual a $\varphi + 1$. Isso é característico da aritmética das relações em geral: se μ e ν são dois números de relação, a regra geral é que $\mu + \nu$ não é igual a $\nu + \mu$. O caso dos ordinais finitos, em que há essa igualdade, é muito excepcional.

A série que finalmente atingimos há pouco consistia, primeiro, em todos os números ímpares e depois em todos os números pares, e seu membro serial é 2φ . Esse número é do maior que φ ou $\varphi + n$, onde n é finito. Convém observar que, em conformidade com a definição geral de ordem, cada um desses arranjos de números inteiros deve ser visto como resultando de alguma relação definida. Por exemplo, aquela que meramente remove 2 para o fim será definida pela seguinte relação: “ x e y são números inteiros finitos, e ou y é 2 e x não é 2, ou nenhum deles é 2 e x é menor do que y ”. Aquela série que põe primeiro todos os números ímpares e depois todos os números pares será definida por: “ x e y são

números inteiros finitos, e ou x é ímpar e y é par ou x é menor do que y e ambos são ímpares ou ambos são pares”. Doravante, via de regra, não nos daremos ao trabalho de dar essas fórmulas; mas o fato de que elas *podiam* ser dadas é essencial.

O número que chamamos 2φ , a saber, o número de uma série que consiste em duas progressões, é por vezes chamado $\varphi \cdot 2$. A multiplicação, como a adição, depende da ordem dos fatores: uma progressão de pares dá uma série do tipo

$$x_1, y_1, x_2, y_2, x_3, y_3, \dots, x_n, y_n, \dots,$$

que é ela própria uma progressão; mas um par de progressões dá uma série que é duas vezes mais longa que uma progressão. É necessário, portanto, distinguir entre 2φ e $\varphi \cdot 2$. O uso é variável; usaremos 2φ para um par de progressões e $\varphi \cdot 2$ para uma progressão de pares, e essa decisão, é claro, governará nossa interpretação geral de “ $\beta\beta$ ” quando φ e β forem números de relação: “ $\varphi \cdot \beta$ ” terão de representar uma soma adequadamente construída de relações α tendo cada uma termos β .

Podemos levar indefinidamente adiante o processo de refinar os números indutivos. Por exemplo, podemos pôr primeiro os números ímpares, depois seus duplos, depois os duplos destes, e assim por diante. Obtemos dessa maneira a série:

$$1, 3, 5, 7, \dots; 2, 6, 10, 14, \dots; 4, 12, 20, 28, \dots; 8, 24, 40, 56, \dots,$$

cujos números são φ^2 , visto que é uma progressão de progressões. Qualquer uma das progressões nessa nova série pode, é claro, ser refinada tal como refinamos nossa progressão original. Podemos prosseguir para φ^3 , φ^4 , $\dots$ φ^4 , e assim por diante; por mais longe que tenhamos ido, sempre podemos ir além.

A série de todos os ordinais que pode ser obtida desta maneira, isto é, tudo o que se pode obter refinando uma progressão, é ela mesma mais longa que qualquer série que pode ser obtida pelo rearranjo dos termos de uma progressão. (Não é difícil provar isso.) Pode-se mostrar que o número cardinal das classes desses ordinais é maior que n_0 ; trata-se do número que Cantor chama n_1 . O número ordinal da série de todos os ordinais que pode ser formada a partir de um n_0 , tomados em ordem de magnitude, é

chamado n_1 . Portanto, uma série cujo número ordinal é φ tem um campo cujo número cardinal é n_1 .

Podemos prosseguir de φ_1 e n_1 para φ_2 e n_2 por um processo exatamente análogo àquele pelo qual avançamos de φ e n_0 para φ_1 e n_1 . E nada nos impede de avançar indefinidamente desta maneira para novos cardinais e novos ordinais. Não se sabe se 2^{n_0} é igual a algum dos cardinais na série dos álefos. Não se sabe sequer se é comparável a eles em magnitude; por tudo que sabemos, não poderia ser nem igual a nenhum dos álefos, nem maior nem menor do que eles. Essa questão está vinculada com o axioma multiplicativo, de que trataremos mais tarde.

Todas as séries que consideramos até agora neste capítulo foram o que se chama “bem ordenadas”. Uma série bem ordenada é aquela que tem um início, tem termos consecutivos e tem um termo *seguinte* após qualquer seleção de seus termos, contanto que haja quaisquer termos após a seleção. Isso exclui, por um lado, séries compactas, em que há termos entre quaisquer dois, e, por outro lado, séries que não têm início ou em que há partes subordinadas sem nenhum início. A série dos números inteiros negativos em ordem de magnitude, como não tem um início mas termina com -1 , não é bem-ordenada; mas tomada na ordem inversa, começando com -1 , é bem-ordenada, sendo de fato uma progressão. A definição é: uma série “bem-ordenada” é aquela em que todas as subclasses (exceto, é claro, a classe nula) têm um primeiro termo.

Um número “ordinal” significa o número de relação de uma série bem ordenada. É, portanto, uma espécie de número serial.

Entre séries bem-ordenadas, aplica-se uma forma generalizada de indução matemática. Pode-se dizer que uma propriedade é “transfinitamente hereditária” se, quando ela pertence a certa seleção dos termos numa série, pertence também ao sucessor imediato deles, contanto que eles tenham um. Numa série bem-ordenada, uma propriedade transfinitamente hereditária pertencente ao primeiro termo da série pertence a toda a série. Isso torna possível provar muitas proposições concernente a séries bem-ordenadas que não são verdadeiras com relação a todas as séries.

É fácil arranjar os números indutivos em séries que não sejam bem-ordenadas, e até arranjá-los em séries compactas. Por exemplo, podemos adotar o seguinte plano: consideremos os decimais de $\cdot 1$ (inclusive) a 1 (exclusive), arranjados em ordem de magnitude. Eles formam uma ordem

compacta; entre quaisquer dois há sempre um número infinito de outros. Se omitirmos o ponto no início de cada uma, teremos uma série compacta que consistirá em todos os números inteiros finitos exceto os divisíveis por 10. Se desejarmos incluir os divisíveis por 10, não há dificuldade; em vez de começar com $\cdot 1$, incluiremos todos os decimais menores do que 1, mas quando removermos o ponto, transferiremos para a direita quaisquer 0s que ocorram no início de nossa decimal. Omitindo-os, e retornando àqueles que não têm nenhum 0 no início, podemos formular a regra para o arranjo de nossos números inteiros da seguinte maneira: de dois números inteiros que não comecem com o mesmo dígito, aquele que começa com o dígito menor vem primeiro. De dois que começam com o mesmo dígito, mas diferem no segundo dígito, aquele que tem o menor segundo dígito vem primeiro; antes de todos, porém, vem aquele que não tem segundo dígito; e assim por diante. Em geral, se dois números inteiros que concordam no tocante aos primeiros n dígitos, mas não com relação ao $(n+1)$ º, vem primeiro aquele que ou não tem nenhum $(n+1)$ º dígito, ou tem um menor que o outro. Essa regra de arranjo dá origem a uma série compacta que contém todos os números inteiros não divisíveis por 10; e, como vimos, não há dificuldade em incluir os divisíveis por 10. Segue-se desse exemplo que é possível construir séries compactas com n_0 termos. De fato, já vimos que há n_0 razões, e razões em ordem de magnitude formam uma série compacta; assim, temos aqui mais um exemplo. Resumiremos esse tópico no próximo capítulo.

Os cardinais transfinitos obedecem a todas as leis formais usuais da adição, multiplicação e exponenciação, mas os ordinais transfinitos só obedecem a algumas delas, e aquelas que são obedecidas por eles são obedecidas por todos os números de relação. Por “leis formais usuais” entendemos as seguintes:

I. A lei comutativa:

$$\alpha + \beta = \beta + \alpha \quad \text{e} \quad \alpha \leftrightarrow \beta = \beta \leftrightarrow \alpha.$$

II. A lei associativa:

$$(\alpha + \beta) + \gamma = \alpha + (\beta + \gamma) \quad \text{e} \quad (\alpha \leftrightarrow \beta) \leftrightarrow \gamma = \alpha \leftrightarrow (\beta \leftrightarrow \gamma).$$

III. A lei distributiva

$$\alpha(\beta + \gamma) = \alpha\beta + \alpha\gamma.$$

Quando a lei comutativa não vigora, a forma demonstrada da lei distributiva deve ser distinguida de

$$(\beta + \gamma) \varphi = \beta x + \gamma x .$$

Como veremos imediatamente, uma forma pode ser verdadeira, e a outra, falsa.

IV. As leis da exponenciação:

$$\alpha^\beta \cdot \alpha^\gamma = \alpha^{\beta+\gamma}, \alpha^\gamma \cdot \beta^\gamma = (\alpha\beta)^\gamma, (\alpha^\beta)^\gamma = \alpha^{\beta\gamma}.$$

Todas essas leis vigoram para cardinais, sejam finitos ou infinitos, e para ordinais *finitos*. Mas quando chegamos aos ordinais infinitos, ou de fato a números de relação em geral, algumas vigoram e outras não. A lei comutativa não vigora; a lei associativa vigora; a lei distributiva (adotando-se a convenção que adotamos previamente com relação à ordem dos fatores num produto) vigora na forma

$$(\beta + \gamma) \alpha = \beta \alpha + \gamma \alpha,$$

mas não na forma

$$\alpha(\beta + \gamma) = \alpha\beta + \alpha\gamma;$$

as leis exponenciais

$$\alpha^\beta \cdot \alpha^\gamma = \alpha^{\beta+\gamma} \text{ e } (\alpha^\beta)^\gamma = \alpha^{\beta\gamma}$$

ainda vigoram, mas não a lei

$$a^y \cdot b^y = (ab)^y,$$

que está obviamente relacionada à lei comutativa para a multiplicação.

As definições de multiplicação e exponenciação pressupostas nas proposições anteriores são um pouco complicadas. O leitor que desejar saber quais são elas e como as leis citadas são provadas deve consultar o segundo volume dos *Principia Mathematica*, notas 172-6.

A aritmética transfinita ordinal foi desenvolvida por Cantor num estágio anterior que a aritmética transfinita cardinal, porque há vários usos matemáticos técnicos que o levaram a ela. Do ponto de vista da filosofia da matemática, porém, ela é menos importante e menos fundamental que a teoria dos cardinais transfinitos. Os cardinais são essencialmente mais

simples que os ordinais, e é um acidente histórico curioso que tenham aparecido de início como uma abstração desses últimos, e só gradualmente tenham passado a ser estudados por si mesmos. Isso não se aplica ao trabalho de Frege, em que os cardinais, os finitos e os transfinitos foram tratados de maneira completamente independente dos ordinais; mas foi o trabalho de Cantor que tornou o mundo ciente do assunto, ao passo que Frege continuou quase desconhecido, provavelmente sobretudo em razão da dificuldade de seu simbolismo. E os matemáticos, como outras pessoas, têm mais dificuldade em compreender e usar noções comparativamente “simples” no sentido lógico do que em manipular noções mais complexas que tenham mais afinidade com sua prática comum. Por essas razões, a verdadeira importância dos cardinais na filosofia matemática só foi reconhecida pouco a pouco. A importância dos ordinais, embora não seja pequena em absoluto, é nitidamente menor do que a dos cardinais, e está em grande parte mesclada à da concepção mais geral de números de relação.

Capítulo 10

Limites e continuidade

Continuamente, foi-se descobrindo que a importância da concepção de “limite” em matemática era maior do que se pensara. A totalidade do cálculo diferencial e integral, na realidade tudo na matemática superior, depende de limites. Outrora, supunha-se que infinitesimais estavam envolvidos nos fundamentos dessas matérias, mas Weierstrass mostrou que isso é um erro: em todos os lugares em que se pensava que ocorriam infinitesimais, o que realmente ocorre é um conjunto de quantidade finitas que têm zero por seu limite inferior. Costumava-se pensar que “limite” era uma noção essencialmente quantitativa, a saber, a noção de uma quantidade da qual outras se aproximavam cada vez mais, de tal modo que entre essas outras havia algumas diferindo por menos que qualquer quantidade designada. De fato, porém, a noção de “limite” é uma noção puramente ordinal, que não envolve quantidade em absoluto (exceto por acidente, quando por acaso a série em questão é quantitativa). Um dado ponto numa linha pode ser o limite de um conjunto de pontos na linha, sem que seja necessário introduzir coordenadas ou mensuração ou nada de quantitativo. O número cardinal n_0 é o limite (na ordem de magnitude) dos números cardinais $1, 2, 3, \dots, n, \dots$, embora a diferença numérica entre n_0 e um cardinal finito seja constante e infinita: de um ponto de vista quantitativo, números finitos não ficam em nada mais próximos de n_0 à medida que aumentam. O que torna n_0 o limite dos números finitos é o fato de que, na série, ele vem imediatamente depois deles, o que é um fato *ordinal*, não um fato quantitativo.

Há várias formas da noção de “limite”, de crescente complexidade. A forma mais simples e mais fundamental, da qual as demais são derivadas, já foi definida, mas repetiremos aqui as definições que conduzem a ela, numa forma geral em que não exigem que a relação envolvida seja serial.

As definições são as seguintes. Os “*mínimos*” de uma classe α com respeito a uma relação P são aqueles membros de α e o campo de P (se houver) com que nenhum membro de α tem a relação P . Os “*máximos*” com respeito a P são os *mínimos* com respeito ao inverso de P . Os “*seqüentes*” de uma classe α com respeito a uma relação P são os *mínimos* dos “*sucessores*” de α , e os “*sucessores*” de α são aqueles membros do campo de P com que todos os membros da parte comum de α e o campo de P tem a relação P . Os “*precedentes*” com respeito a P são os seqüentes com respeito ao inverso de P . Os “*limites superiores*” de α com respeito a P são os seqüentes contanto que α não tenha nenhum máximo; mas se α tiver um máximo, não tem nenhum limite superior. Os “*limites inferiores*” com respeito a P são os limites superiores com respeito ao inverso de P .

Sempre que P tiver conexidade, uma classe não pode ter mais que um máximo, um mínimo, um seqüente etc. Assim, nos casos em que estamos interessados na prática, podemos falar de “o limite” (se houver algum).

Quando P for uma relação serial, podemos simplificar muito essa definição de limite. Podemos, nesse caso, definir primeiro a “*fronteira*” de uma classe α , isto é, seus limites ou máximos, e depois avançar para distinguir o caso em que a fronteira é o limite daquele em que ela é um máximo. Para essa finalidade, é melhor usar a noção de “segmento”.

Falaremos do “segmento de P definido por uma classe α ” como todos aqueles termos que têm a relação P como um ou mais membros de α . Esse será um segmento no sentido definido no Capítulo 7; na realidade, todo segmento no sentido ali definido é o segmento definido por *alguma* classe α . Se P for serial, o segmento definido por α consistirá em todos os termos que precedem um termo ou outro de α . Se α tiver um máximo, o segmento será todos os predecessores do máximo. Mas se α não tiver máximo, todo membro de α precederá algum outro membro de α , e a totalidade de α estará portanto incluída no segmento definido por α . Tomemos, por exemplo, a classe que consiste nas frações

$$1/2, 3/4, 7/8, 15/16, \dots,$$

isto é, em todas as frações da forma $1 - 1/2^n$ para diferentes valores finitos de n . Esta série de frações não tem máximo, e é claro que o segmento que ela define (em toda a série de frações em ordem de magnitude) é a classe de todas as frações próprias. Ou ainda, tomemos os números primos,

considerados uma seleção dos cardinais (finitos e infinitos) em ordem de magnitude. Nesse caso o segmento definido consiste em todos os números inteiros finitos.

Supondo que P é serial, a “fronteira” de uma classe α será o termo x (se ele existir) cujos predecessores são o segmento definido por α .

Um “máximo” de α é uma fronteira que é membro de α .

Um “limite superior” de α é uma fronteira que não é membro de α .

Se uma classe não tiver fronteira, não terá máximo nem limite. Esse é o caso de um corte de Dedekind “irracional”, ou do que é chamado “lacuna”.

Assim o “limite superior” de um conjunto de termos α com respeito a uma série P é aquele termo x (se ele existir) que vem depois de todos os α 's, mas é tal que todo termo anterior vem antes de alguns dos α 's.

Podemos definir todos os “pontos limitantes superiores” de um conjunto de termos β como todos aqueles que são os limites superiores de conjuntos de termos escolhidos de β . Teremos, é claro, de distinguir pontos limitantes superiores de pontos limitantes inferiores. Se considerarmos, por exemplo, a série de números ordinais:

$$1, 2, 3, \dots \varphi, \varphi + 1, \dots 2\varphi, 2\varphi + 1, \dots 3\varphi, \dots \varphi^2, \dots \varphi^3, \dots,$$

os pontos limitantes superiores do campo dessa série são aqueles que não têm predecessores imediatos, isto é,

$$1, \varphi, 2\varphi, 3\varphi, \dots \varphi^2, \varphi^2 + \varphi, \dots 2\varphi^2, \dots \varphi^3 \dots$$

Os pontos limitantes superiores do campo dessa nova série serão

$$1, \varphi^2, 2\varphi^2, \dots \varphi^3, \varphi^3 + \varphi^2$$

Por outro lado, a série dos ordinais — na verdade, toda série bem-ordenada — não tem pontos limitantes inferiores, porque não há termos, exceto o último, que não tenham sucessores imediatos. Mas se considerarmos tal série a série das razões, cada membro dela um ponto limitante tanto superior quanto inferior para conjuntos adequadamente escolhidos. Se considerarmos a série dos números reais, e selecionarmos a partir dela os números reais racionais, esse conjunto (os racionais) terá todos os números reais como pontos limitantes superiores e inferiores. Os pontos limitantes de um conjunto são chamados sua “primeira derivada”;

os pontos limitantes da primeira derivada são chamados a segunda derivada, e assim por diante.

Com relação a limites, podemos distinguir vários graus do que pode ser chamado “continuidade” numa série. A palavra “continuidade” havia sido usada por muito tempo, mas permanecera sem nenhuma definição precisa até o tempo de Dedekind e Cantor. Cada um desses dois homens deu um significado preciso ao termo, mas a definição de Cantor é mais estreita que a de Dedekind: uma série que tem a continuidade cantoriana deve ter a continuidade dedekindiana, mas a recíproca não é verdadeira.

A primeira definição que ocorreria naturalmente a um homem que buscasse um sentido preciso para a continuidade de séries seria defini-la como consistindo no que chamamos “compacidade”, isto é, no fato de haver outros termos entre quaisquer dois termos da série. Mas isso seria uma definição inadequada, por causa da existência de “lacunas” em séries como a série das razões. Vimos no Capítulo 7 que há inúmeras maneiras pelas quais a série das razões pode ser dividida em duas partes, das quais uma precede inteiramente a outra, e das quais a primeira não tem último termo, ao passo que a segunda não tem primeiro termo. Tal estado de coisas parece contrário à vaga impressão que temos quanto ao que deveria caracterizar a “continuidade”, e, mais ainda, mostra que a série das razões não é o tipo de série necessária para muitas finalidades matemáticas. Tomemos a geometria, por exemplo: desejamos poder dizer que, quando duas linhas retas se cruzam, elas têm um ponto em comum, mas se a série de pontos numa linha fosse similar à série das razões, as duas linhas poderiam se cruzar numa “lacuna” e não ter nenhum ponto em comum. Esse é um exemplo grosseiro, mas seria possível dar muitos outros para mostrar que a compacidade é inadequada como definição matemática de continuidade.

Foram as necessidades da geometria, mais do que qualquer outra coisa, que levaram à definição da continuidade “dedekindiana”. Lembremos que definimos uma série como dedekindiana quando todas as subclasses do campo têm uma fronteira. (É suficiente dar por certo que há sempre uma fronteira *superior*, ou que há sempre uma fronteira *inferior*. Se damos uma dessas por certo, a outra pode ser deduzida.) Isso quer dizer que uma série é dedekindiana quando não há nenhuma lacuna. A ausência de lacunas pode resultar ou do fato de os termos terem sucessores, ou da existência de limites na ausência de máximos. Assim, uma série finita ou uma série

bem-ordenada é dedekindiana, e o mesmo pode ser dito da série dos números reais. O primeiro tipo de série dedekindiana é excluído admitindo-se que nossa série é compacta; nesse caso, nossa série deve ter uma propriedade que pode, para muitos propósitos, ser adequadamente chamada de continuidade. Somos assim levados à definição: uma série tem “continuidade dedekindiana” quando é dedekindiana e compacta.

Mas esta definição ainda é larga demais para muitas finalidades. Suponhamos, por exemplo, que desejamos poder atribuir tais propriedades a um espaço geométrico de maneira a assegurar que cada ponto possa ser especificado por meio de coordenadas que sejam números reais: isto não é assegurado pela continuidade dedekindiana apenas. Queremos ter certeza de que cada ponto que não possa ser especificado por coordenadas *racionais* possa ser especificado como o limite de uma *progressão* de pontos cujas coordenadas sejam racionais, e esta é uma propriedade adicional que nossa definição não nos permite deduzir.

Somos assim levados a uma investigação mais atenta das séries com respeito a limites. Essa investigação foi feita por Cantor e formou a base de sua definição de continuidade, embora, em sua forma mais simples, a definição oculte de certo modo as considerações que lhe deram origem. Por isso, antes de dar a definição de continuidade de Cantor, repassaremos algumas de suas concepções sobre este assunto.

Cantor define uma série como “perfeita” quando todos os seus pontos limitantes e todos os seus pontos limitantes pertencem a ela. Mas essa definição não expressa com muita precisão o que ele quer dizer. Não há necessidade de nenhuma correção no que concerne à propriedade de todos os pontos da série serem pontos limitantes; essa é uma propriedade que pertence às séries compactas e a nenhum outro tipo de série, se for preciso que todos os pontos sejam limitantes superiores ou limitantes inferiores. Mas se supomos apenas que eles são pontos limitantes de uma maneira, sem especificar qual, outras séries terão a propriedade em questão — por exemplo, a série dos decimais em que um decimal que termine num 9 recorrente é distinguida do decimal terminante correspondente e posto imediatamente antes dele. Tal série aproxima-se muito de ser compacta, mas tem termos excepcionais que são consecutivos, o primeiro dos quais não tem predecessor imediato, ao passo que o segundo não tem sucessor imediato. Afora séries desse tipo, as séries em que cada ponto é um ponto limitante são compactas; e isso é

válido sem qualificação, se for especificado que cada ponto deve ser um ponto limitante superior (ou que cada ponto deve ser um ponto limitante inferior).

Embora Cantor não considere explicitamente a matéria, devemos distinguir diferentes tipos de pontos limitantes segundo a natureza da menor subsérie pela qual eles podem ser definidos. Cantor supõe que eles devem ser definidos por progressões, ou por regressões (que são o inverso de progressões). Quando todo membro de nossa série é o limite de uma progressão ou regressão, Cantor chama nossa série “condensada em si mesma” (*insichdicht*).

Chegamos agora à segunda propriedade pela qual a perfeição devia ser definida, a saber, a propriedade que Cantor chama de ser “fechada” (*abgeschlossen*). Isso, como vimos, foi definido primeiro como consistindo no fato de todos os pontos limitantes de uma série pertencerem a ela. Mas isso só tem um significado efetivo se nossa série for *dada* como contida em alguma outra série maior (no caso, por exemplo, da seleção de números reais), e pontos limitantes forem tomados em relação a séries maiores. De outro modo, se for considerada simplesmente em si mesma, uma série não pode deixar de conter seus pontos limitantes. O que Cantor *quer dizer* não é exatamente o que ele diz; na realidade, em outras ocasiões diz algo bastante diferente, que é o que quer dizer. O que realmente quer dizer é que toda série subordinada que se esperaria que tivesse um limite tem de fato um limite dentro da série dada; isto é, toda série subordinada que não tem máximo tem um limite, isto é, toda série subordinada tem uma fronteira. Mas Cantor não afirma isto com relação a toda série subordinada, mas apenas com relação a progressões e regressões. (Não fica claro até que ponto ele reconhece que isso é uma limitação.) Assim, finalmente, vemos que a definição que queremos é a seguinte: diz-se que uma série é “fechada” (*abgeschlossen*) quando todas as progressões ou regressões contidas na série têm um limite na série.

Temos depois a definição adicional: uma série é “perfeita” quando é *condensada em si mesma* e *fechada*, isto é, quando cada termo é o limite de uma progressão ou regressão, e todas as progressões ou regressões contidas na série têm um limite na série.

Ao procurar uma definição de continuidade, o que Cantor tem em mente é a busca de uma definição que se aplique à série dos números reais e a qualquer série similar a essa, mas não a outras. Para essa finalidade, temos

de acrescentar mais uma propriedade. Entre os números reais, alguns são racionais, alguns são irracionais; embora o número de irracionais seja maior do que o número de racionais, há racionais entre quaisquer dois números reais, por menos diferentes que sejam. O número de racionais, como vimos, é n_0 . Isto dá uma propriedade adicional que é suficiente para caracterizar completamente a continuidade, a saber, a propriedade de conter uma classe de n_0 membros de tal maneira que alguns dessa classe ocorram entre dois termos quaisquer de nossa série, por mais próximos que sejam. Essa propriedade, acrescentada à perfeição, é suficiente para definir uma classe de séries que são todas similares e são de fato um número serial. Cantor define essa classe como a das séries contínuas.

Podemos simplificar ligeiramente sua definição. Para começar, dizemos:

Uma “classe mediana” de uma série é uma subclasse do campo tal que membros dela podem ser encontrados entre quaisquer dois termos da série.

Assim, os racionais são uma classe mediana na série dos números reais. É evidente que não pode haver classes medianas exceto em séries compactas.

Verificamos então que a definição de Cantor é equivalente à seguinte:

Uma série é “contínua” quando (1) é dedekindiana, (2) contém uma classe mediana com n_0 termos.

Para evitar confusão, falaremos desse tipo como “continuidade cantoriana”. Veremos que ela implica a continuidade dedekindiana, mas a recíproca não é verdadeira. Todas as séries que têm continuidade cantoriana são similares, mas não todas as séries que têm continuidade dedekindiana.

As noções de *limite* e *continuidade* que estivemos definindo não devem ser confundidas com as noções do limite de uma função para aproximações a um dado argumento, ou da continuidade de uma função na vizinhança de um dado argumento. Essas são noções diferentes, muito importantes, mas derivadas das formuladas previamente e mais complicadas. A continuidade do movimento (se o movimento for contínuo) é um caso da continuidade de uma função; por outro lado, a continuidade do espaço e tempo (se forem contínuos) é um caso da continuidade de séries, ou (para falar de maneira mais cautelosa) de um tipo de continuidade que pode, por manipulação matemática suficiente, ser reduzida à continuidade de séries. Em vista da importância fundamental

do movimento na matemática aplicada, bem como por outras razões, convém tratar brevemente das noções de limite e continuidade tal como aplicadas a funções; mas é melhor reservar esse tema para um capítulo separado.

As definições de continuidade que estivemos considerando, a saber, as de Dedekind e Cantor, não correspondem muito estreitamente à vaga idéia que está associada à palavra na mente do homem comum ou do filósofo. Eles concebem a continuidade antes como uma ausência de separação, o tipo de obliteração geral de distinções que caracteriza um nevoeiro denso. Um nevoeiro dá uma impressão de vastidão sem multiplicidade ou divisão definidas. É esse tipo de coisa que um metafísico entende por “continuidade”, declarando, com toda razão, que ela é característica de sua vida mental e da vida das crianças e dos animais.

A idéia geral vagamente indicada pela palavra “continuidade” quando empregada dessa maneira, ou pela palavra “fluxo”, é certamente muito diferente da que estivemos definindo. Tomemos, por exemplo, a série dos números reais. Cada um é o que é, de maneira muito definida e cabal; não se converte no outro por graus imperceptíveis; é uma unidade permanente, separada, e a distância que o separa de qualquer outra unidade é finita, embora possa se tornar menor do que qualquer quantidade finita dada designada de antemão. A questão da relação entre o tipo de continuidade existente em meio aos números reais e o tipo exibido, por exemplo, pelo que vemos num dado momento, é difícil e intrincada. Não se pode sustentar que os dois tipos são idênticos, mas é perfeitamente possível, penso eu, sustentar que a concepção matemática que consideramos nesse capítulo fornece o esquema lógico abstrato a que é possível referir material empírico por manipulação adequada, se esse material puder ser chamado “contínuo” em algum sentido precisamente definível. Seria inteiramente impossível justificar essa tese dentro dos limites do presente volume. O leitor interessado poderá ler uma tentativa de justificá-la, no tocante ao *tempo* em particular, pelo presente autor no *Monist* para 1914-15, bem como em partes de *Our Knowledge of the External World*. Com essas indicações, devemos deixar esse problema, por mais interessante que seja, para retornar a tópicos mais estreitamente vinculados à matemática.

Capítulo 11

Limites e continuidade de funções

Neste capítulo trataremos da definição do limite de uma função (se houver um) à medida que o argumento se aproxima de um valor dado, e também da definição do que se quer dizer por uma “função contínua”. Essas idéias são ambas um pouco técnicas e dificilmente precisariam ser abordadas numa mera introdução à filosofia matemática, a não ser pelo fato de que, especialmente por meio do chamado cálculo infinitesimal, idéias erradas sobre nossos tópicos presentes tornaram-se tão firmemente arraigadas nas mentes de filósofos profissionais que se faz necessário um esforço prolongado e considerável para erradicá-las. Pensa-se, desde o tempo de Leibniz, que o cálculo diferencial e integral requer quantidades infinitesimais. Os matemáticos (em especial Weierstrass) provaram que isso é um erro; mas erros incorporados, *e.g.*, no que Hegel tem a dizer sobre matemática, dificilmente morrem, e a tendência é que os filósofos a ignorem o trabalho de homens como Weierstrass.

Limites e continuidade de funções, em obras sobre matemática comum, são definidos em termos que envolvem número. Isso não é essencial, como o mostrou o dr. Whitehead.¹ Começaremos, contudo, com as definições que aparecem nos livros-texto e depois prosseguiremos para mostrar como elas podem ser generalizadas de maneira a se aplicarem a séries em geral, e não somente às numéricas ou numericamente mensuráveis.

Consideremos uma função matemática ordinária fx , onde x e fx são ambos números reais, e fx tem um único valor — isto é, quando x é dado, há somente um valor que fx pode ter. Chamamos x o “argumento” e fx o “valor para o argumento x ”. Quando uma função é o que chamamos “contínua”, a idéia aproximada para a qual estamos procurando uma definição precisa é que pequenas diferenças em x devem corresponder a pequenas diferenças em fx , e se tornarmos as diferenças em x pequenas o bastante, poderemos fazer as diferenças em fx caírem abaixo de qualquer

quantidade designada. Não queremos, se uma função deve ser contínua, que haja saltos súbitos, de modo que, para algum valor de x , alguma mudança, por menor que seja, provoque uma mudança em fx que exceda alguma quantidade finita designada. As funções simples ordinárias da matemática têm esta propriedade: ela pertence, por exemplo, a x^2 , x^3 , $\dots$ $\log x$, $\sin x$, e assim por diante. Mas não é difícil definir funções descontínuas. Tome, como um exemplo não-matemático, “o lugar de nascimento da pessoa mais jovem que vive no momento t ”. Isso é uma função de t ; seu valor é constante desde o momento do nascimento de uma pessoa até o momento do nascimento seguinte, e então o valor muda *subitamente* de um lugar de nascimento para outro. Um exemplo matemático análogo seria “o mais próximo número inteiro abaixo de x ”, onde x é um número real. Essa função permanece constante de um número inteiro para o seguinte, e então dá um salto súbito. A realidade é que, embora sejam mais conhecidas, as funções contínuas são exceções: há infinitamente mais funções descontínuas que contínuas.

Muitas funções são descontínuas para um ou diversos valores da variável, mas contínuas para todos os outros. Tomemos como exemplo $\sin 1/x$. A função $\sin \theta$ passa por todos os valores de -1 a 1 cada vez que θ passa de $-\pi/2$ para $\pi/2$, ou de $\pi/2$ para $3\pi/2$, ou em geral de $(2n - 1)\pi/2$ para $(2n + 1)\pi/2$, onde n é qualquer número inteiro. Mas se considerarmos $1/x$ quando x é muito pequeno, vemos que, à medida que x diminui, $1/x$ cresce cada vez mais depressa, de tal modo que passa cada vez mais rapidamente pelo ciclo de valores de um múltiplo de $\pi/2$ para outro à medida que x se torna cada vez menor. Em consequência, $\sin 1/x$ passa cada vez mais rapidamente de -1 para 1 e de novo para -1 à medida que x diminui. De fato, se tomarmos qualquer intervalo contendo 0 , digamos o intervalo de $-\varepsilon$ para $+\varepsilon$ onde ε é algum número muito pequeno, $\sin 1/x$ passará por um número infinito de oscilações nesse intervalo, e não podemos diminuir as oscilações tornando o intervalo menor. Assim, nas vizinhanças do argumento 0 a função é descontínua. É fácil produzir funções que sejam descontínuas em vários lugares, ou em n_0 lugares, ou em todos os lugares. Exemplos serão encontrados em qualquer livro sobre a teoria das funções de uma variável real.

Passando agora a procurar uma definição precisa do que se tem em mente ao dizer que uma função é contínua para um dado argumento, quando argumento e valor são ambos números reais, definamos primeiro a

“vizinhança” de um número x como todos os números de $x - \varepsilon$ a $x + \varepsilon$, em que ε é algum número que, em casos importantes, será muito pequeno. É claro que a continuidade num ponto dado tem a ver com o que acontece em *qualquer* vizinhança desse ponto, por menor que seja.

O que desejamos é isto: se a for o argumento para o qual desejamos que nossa função seja contínua, definamos primeiro uma vizinhança (digamos α) contendo o valor fa que a função tem para o argumento a ; desejamos que, se tomarmos uma vizinhança suficientemente pequena contendo a , todos os valores para argumentos de um extremo ao outro dessa vizinhança deverão estar contidos na vizinhança α , por menor que tenhamos feito α . Isso quer dizer que, se decretamos que nossa função não deve diferir de fa por mais de uma quantidade extremamente minúscula, podemos sempre encontrar uma extensão de números reais, tendo a no meio dela, tal que de um extremo a outro dessa extensão fx não difira de fa por mais do que uma quantidade minúscula estabelecida. E isso deve permanecer verdadeiro seja qual for a quantidade minúscula que possamos selecionar. Somos levados portanto à seguinte definição: Diz-se que a função $f(x)$ é “contínua” para o argumento a se, para cada número positivo σ , diferente de 0, mas tão pequeno quanto desejemos, existir um número positivo ε , diferente de 0, tal que, para todos os valores de δ numericamente menores do que² ε , a diferença $f(a+\delta) - f(a)$ seja numericamente menor do que σ .

Até agora não definimos o “limite” de uma função para um dado argumento. Se o tivéssemos feito, teríamos podido definir a continuidade de uma função de outra maneira: uma função é contínua num ponto em que seu valor é o mesmo que o limite de seu valor para aproximações tanto a partir de cima quanto a partir de baixo. Mas só a função excepcionalmente “dócil” tem um limite definido quando o argumento se aproxima de um dado ponto. A regra geral é que uma função oscile, e que, dada qualquer vizinhança de um dado argumento, por menor que seja, toda uma extensão de valores ocorrerá para argumentos dentro daquela vizinhança. Como essa é a regra geral, vamos considerá-la primeiro.

Consideremos o que pode acontecer quando o argumento se aproxima de algum valor a a partir de baixo. Isto é, queremos considerar o que acontece para argumentos contidos no intervalo de $a - \varepsilon$ a a , onde ε é algum número que, em casos importantes, será muito pequeno.

Os valores da função para argumentos de $a - \varepsilon$ a a (a excluído) serão um conjunto de números reais que definirão uma certa seção do conjunto de números reais, a saber, a seção que consiste naqueles números não maiores do que *todos* os valores para argumentos de $a - \varepsilon$ a a . Dado qualquer número nessa seção, haverá valores pelo menos tão grandes quanto esse número para argumentos entre $a - \varepsilon$ a a , isto é, para argumentos que caem a muito pouca distância de a (se ε for muito pequeno). Tomemos todos os ε 's possíveis e todas as seções correspondentes possíveis. Chamaremos a parte comum de todas essas seções a “seção extrema” à medida em que o argumento se aproxima de a . Dizer que um número z pertence à seção extrema é dizer que, por menor que possamos fazer ε , há argumentos entre $a - \varepsilon$ para os quais o valor da função é não menor do que z .

Podemos aplicar exatamente o mesmo processo a seções superiores, isto é, seções que vão de algum ponto até o topo, em vez de irem da parte inferior a algum ponto. Aqui tomamos aqueles números não *menores* do que todos os valores para argumento de $a - \varepsilon$ a a ; isso define uma seção superior que variará à medida que ε varie. Tomando a parte comum de todas essas seções para todos os ε 's possíveis, obtemos a “seção superior extrema”. Dizer que um número z pertence à seção superior extrema é dizer que, por menor que façamos ε , haverá argumentos entre $a - \varepsilon$ a a para os quais o valor da função não será *maior* do que z .

Se o termo z pertencer tanto à seção extrema quanto à seção superior extrema, diremos que ele pertence à “oscilação extrema”. Podemos ilustrar a matéria considerando mais uma vez a função $\sin 1/x$ à medida que x se aproxima do valor 0. Suporemos, para nos pôr de acordo com as definições apresentadas, que a aproximação a esse valor se dá a partir de baixo.

Começemos com a “seção extrema”. Entre $-\varepsilon$ e 0, qualquer que seja ε , a função assumirá o valor 1 para certos argumentos, mas jamais assumirá um valor maior. Portanto, a seção extrema consiste em todos os números reais, positivos e negativos, até e incluindo 1; isto é, consiste em todos os números negativos juntamente com 0, juntamente com os números positivos até e incluindo 1.

De maneira semelhante, a “seção superior extrema” consiste em todos os números positivos juntamente com 0, juntamente com os números negativos até e incluindo -1 .

Assim a “oscilação extrema” consiste em todos os números reais de -1 a 1 , ambos incluídos.

Podemos dizer de maneira geral que a “oscilação extrema” de uma função à medida que o argumento se aproxima de a a partir de baixo consiste em todos aqueles números x tais que, por mais perto que cheguemos de a , ainda encontraremos valores tão grandes quanto x e valores tão pequenos quanto x .

A oscilação extrema pode não conter nenhum termo, ou conter um termo, ou muitos. Nos dois primeiros casos, a função tem um limite definido para aproximações a partir de baixo. Se a oscilação extrema tiver um termo, isso é bastante óbvio. É igualmente óbvio se não tiver nenhum, pois não é difícil provar que, se a oscilação extrema for nula, a fronteira da seção extrema será a mesma da seção superior extrema e poderá ser definida como o limite da função para aproximações a partir de baixo. Mas se a oscilação extrema tiver muitos termos, não haverá limite definido para a função para aproximações a partir de baixo. Nesse caso, podemos tomar as fronteiras superior e inferior da oscilação extrema (isto é, a fronteira inferior da seção superior extrema e a fronteira superior da seção extrema) como os limites superior e inferior de seus valores “extremos” para aproximações a partir de baixo. De maneira similar, obtemos limites inferior e superior dos valores “extremos” para aproximações a partir de cima. Assim temos, no caso geral, *quatro* limites para uma função para aproximações a um dado argumento. O limite para um dado argumento a só existe quando todos esses quatro são iguais, sendo, portanto, o valor comum deles. Se esse for também o valor para o argumento a , a função será contínua para esse argumento. Isso pode ser tomado como definindo continuidade: é equivalente à nossa definição anterior.

Podemos definir o limite de uma função para um dado argumento (se ele existir) sem passar pela oscilação extrema e os quatro limites do caso geral. A definição origina-se exatamente como a definição anterior de continuidade. Definamos o limite para aproximações a partir de baixo. Para que haja um limite definido para aproximações a a a partir de baixo, é necessário e suficiente que, dado qualquer número pequeno σ , dois valores para argumentos suficientemente próximos de a (mas ambos menores do que a) difiram por menos que σ ; isto é, se ε for suficientemente pequeno, e nossos argumentos se encontrarem ambos

entre $a - \varepsilon$ e a (a excluído), então a diferença entre os valores para esses argumentos será menor do que σ . Isso deve se aplicar a qualquer σ , por menor que seja; nesse caso, a função tem um limite para aproximações a partir de baixo. De maneira semelhante, definimos o caso em que há um limite para aproximações a partir de cima. Esses dois limites, mesmo quando ambos existem, não precisam ser idênticos; e se forem idênticos, ainda assim não precisam ser idênticos ao *valor* para o argumento a . É apenas nesse último caso que chamamos a função *contínua* para o argumento a .

Uma função é chamada “contínua” (sem qualificação) quando é contínua para todos os argumentos.

Outro método ligeiramente diferente de chegar à definição de continuidade é o seguinte.

Digamos que uma função “converge finalmente para uma classe α ” se houver algum número real tal que, para esse argumento e todos os argumentos maiores do que esse, o valor da função for membro da classe α . De maneira semelhante, diremos que uma função “converge para α à medida que o argumento se aproximar de x a partir de baixo” se houver algum argumento y menor do que x tal que de um extremo ao outro do intervalo de y (incluído) a x (excluído) a função tiver valores que são membros de α . Agora podemos dizer que uma função é contínua para o argumento a , para a qual tem o valor fa , se satisfizer quatro condições, a saber:

(1) Dado qualquer número real menor do que fa , a função converge para os sucessores desse número à medida que o argumento se aproxima de a a partir de baixo;

(2) Dado qualquer número real maior do que fa , a função converge para os predecessores desse número à medida que o argumento se aproxima de a a partir de baixo;

(3) e (4) Condições similares para aproximações a a a partir de cima.

As vantagens dessa forma de definição é que ela analisa as condições de continuidade em quatro, derivadas da consideração de argumentos e valores respectivamente maiores ou menores do que o argumento e o valor para os quais a continuidade deve ser definida.

Podemos agora generalizar nossas definições de modo que elas se apliquem a séries que não sejam numericamente conhecidas ou numericamente mensuráveis. Um caso que convém ter em mente é o do

movimento. Há um conto de H.G. Wells que ilustra, a partir do caso do movimento, a diferença entre o limite de uma função para um dado argumento e seu valor para o mesmo argumento. O herói da história, que possuía, sem o saber, o poder de realizar seus desejos, estava sendo atacado por um policial, mas ao exclamar “Vá à —”, constatou que o policial desapareceu. Se $f(t)$ era a posição do policial no tempo t , e t_0 o momento da exclamação, o limite das posições do policial à medida que t se aproximava de t_0 a partir de baixo estaria em contato com o herói, ao passo que o valor para o argumento t_0 era —. Mas ocorrências desse tipo são supostamente raras no mundo real, e presume-se, embora sem provas adequadas, que todos os movimentos são contínuos, isto é, que, dado qualquer corpo, se $f(t)$ for sua posição no tempo t , $f(t)$ será uma função contínua de t . É o significado de “continuidade” envolvido nessas afirmações que desejamos agora definir de maneira tão simples quanto possível.

As definições dadas para o caso de funções em que o argumento e o valor são números reais podem ser facilmente adaptadas para uso mais geral.

Suponhamos que P e Q são duas relações, que imaginamos serem seriais, embora isso não seja necessário às nossas definições. Suponhamos que R é uma relação um-muitos cujo domínio está contido no campo de P , enquanto seu domínio inverso está contido no campo de Q . Então R é (num sentido generalizado) uma função cujos argumentos pertencem ao campo de Q , enquanto seus valores pertencem ao campo de P . Suponhamos, por exemplo, que estamos tratando de uma partícula que se move numa linha: Q é a série temporal; P , a série de pontos em nossa linha da esquerda para a direita; R , a relação da posição de nossa partícula na linha no instante a com o instante a , de tal modo que “o R de a ” seja sua posição no instante a . Podemos ter essa ilustração em mente por todas as nossas definições.

Diremos que a função R é contínua para o argumento a se, dado qualquer intervalo α na série P contendo o valor da função para o argumento a , houver um intervalo na série Q contendo a não como um ponto final e tal que, de um extremo a outro desse intervalo, a função tenha valores que sejam membros de α . (Por “intervalo”, queremos dizer todos os termos entre quaisquer dois; isto é, se x e y forem dois membros do campo P , e x tiver a relação P com y , entenderemos pelo “intervalo P x

a y ” todos os termos z tais que x tenha a relação P com y — juntamente, quando assim declarado, com x ou y eles próprios.)

Podemos definir facilmente a “seção extrema” e a “oscilação extrema”. Para definir a “seção extrema” para aproximações ao argumento a a partir de baixo, tomemos qualquer argumento y que preceda a (isto é, tenha a relação Q com a), tomemos os valores da função para todos os argumentos até e incluindo y , e formemos a seção de P definida por esses valores, isto é, aqueles membros da série P anteriores a alguns desses valores ou idênticos a eles. Formemos todas aquelas seções para todos os y 's que precedam a , e tomemos sua parte comum; essa será a seção extrema. A seção superior extrema e a oscilação extrema são então definidas exatamente como no caso anterior.

A adaptação da definição de convergência e a definição alternativa resultante de continuidade não oferecem nenhum tipo de dificuldade.

Dizemos que uma função R é “ultimamente Q -convergente para α ” se houver um membro y do domínio inverso de R e do campo de Q tal que o valor da função para o argumento y e para qualquer argumento com que y tenha a relação Q for membro de α . Diremos que R “ Q -converge para α à medida que o argumento se aproximar de um argumento dado a ” se houver um termo y que tenha a relação Q com a e pertença ao domínio inverso de R e tal que o valor da função para qualquer argumento no intervalo Q de y (inclusive) a a (exclusive) pertença a α .

Das quatro condições que uma função deve preencher para ser contínua para o argumento a , a primeira é, tomando b como o valor do argumento a :

Dado qualquer termo que tenha a relação P com b , R Q -converge para os sucessores de b (com respeito a P) à medida que o argumento se aproximar de a a partir de baixo.

Obtemos a segunda condição substituindo P por seu inverso; a terceira e a quarta são obtidas a partir da primeira e da segunda, substituindo-se Q por seu inverso.

Não há nada, portanto, nas noções do limite de uma função ou da continuidade de uma função que envolva essencialmente número. Ambos podem ser definidos de maneira geral, e muitas proposições sobre eles podem ser provadas para quaisquer duas séries (uma sendo a série dos argumentos, e a outra, a série dos valores). Como se vê, as definições não envolvem infinitesimais. Envolvem classes infinitas de intervalos, que

diminuem sem nenhum limite menor do que zero, mas não envolvem quaisquer intervalos que não sejam finitos. Isso é análogo ao fato de que, se uma linha de um centímetro for dividida ao meio, depois novamente dividida, de maneira indefinida, nunca chegaremos a infinitesimais: após n bissecções, o comprimento de nosso pedacinho será $1/2^n$ de um centímetro; e isso é finito, seja qual for o número finito que n possa ser. O processo de sucessivas bissecções não conduz a divisões cujo número ordinal seja infinito, uma vez que é essencialmente um processo um-por-um. Portanto, não se chegará a infinitesimais dessa maneira. A confusão acerca desses tópicos teve muito a ver com as dificuldades encontradas na discussão da infinidade e da continuidade.

Capítulo 12

Seleções e o axioma multiplicativo

Neste capítulo temos de considerar um axioma que pode ser enunciado, mas não provado, em termos de lógica, e que é conveniente, embora não indispensável, em certas porções da matemática. É conveniente no sentido de que muitas proposições interessantes, que parece natural supor verdadeiras, não podem ser provadas sem a sua ajuda; mas não é indispensável, porque mesmo sem essas proposições as matérias em que eles ocorrem ainda existem, embora numa forma um pouco mutilada.

Antes de enunciar o axioma multiplicativo, devemos explicar a teoria das seleções e a definição de multiplicação quando o número de fatores pode ser infinito.

Ao definir as operações aritméticas, o único procedimento correto é construir uma classe real (ou relação, no caso dos números de relação) que tenha o número requerido de termos. Isso exige por vezes certo grau de engenhosidade, mas é essencial para se provar a existência do número definido. Tomemos, como o exemplo mais simples, o caso da adição. Suponhamos que nos seja dado um número cardinal μ , e uma classe α que tenha μ termos. Como definiremos $\mu + \mu$? Para essa finalidade devemos ter *duas* classes tendo μ termos, e elas não devem se superpor. Podemos construir tais classes a partir de α de várias maneiras, das quais a seguinte talvez seja a mais simples: formemos primeiro todos os pares ordenados cujo primeiro termo seja uma classe consistindo em um único membro de α , e cujo segundo termo seja a classe nula; depois, em segundo lugar, formemos todos os pares ordenados cujo primeiro termo seja a classe nula e cujo segundo termo seja uma classe consistindo em um único membro de α . Essas duas classes de pares não têm nenhum membro em comum, e a soma lógica das duas terá $\mu + \mu$ termos. De maneira exatamente análoga podemos definir $\mu + \nu$, dado que μ é o número de alguma classe α e ν é o número de alguma classe β .

Essas definições, via de regra, são meramente questão de um expediente técnico adequado. Mas no caso da multiplicação, em que o número de fatores pode ser infinito, problemas importantes surgem da definição.

Quando o número de fatores é finito, a multiplicação não oferece problemas. Dadas duas classes α e β , das quais a primeira tem μ termos, e a segunda, ν termos, podemos definir $\mu \times \nu$ como o número de pares ordenados que podem ser formados escolhendo-se o primeiro termo em α e o segundo em β . Observemos que essa definição não requer que α e β não se superponham; permanece adequada até quando α e β são idênticas. Por exemplo, suponhamos que α é a classe cujos membros são x_1, x_2, x_3 . Nesse caso a classe usada para definir o produto $\mu \times \mu$ é a classe de pares:

$$(x_1, x_1), (x_1, x_2), (x_1, x_3); (x_2, x_1), (x_2, x_2), (x_2, x_3); (x_3, x_1), (x_3, x_2), (x_3, x_3).$$

Essa definição permanece aplicável quando μ ou ν ou ambos são infinitos, e pode ser estendida passo a passo a três ou quatro ou qualquer número finito de fatores. Nenhuma dificuldade surge no tocante a essa definição, exceto que ela não pode ser estendida a um número *infinito* de fatores.

O problema da multiplicação quando o número de fatores pode ser infinito surge da seguinte maneira: suponhamos que temos uma classe κ que consiste em classes; suponhamos que o número de termos em cada uma dessas classes seja dado. Como definiremos o produto de todos esses números? Se pudermos formular nossa definição de maneira geral, ela será aplicável quer κ seja finito ou infinito. Deve-se observar que o problema é conseguir tratar do caso quando κ é infinito, não com o caso em que seus membros o são. Se κ não for infinito, o método definido acima é igualmente aplicável, quer seus membros sejam finitos ou infinitos. É do caso em que κ é infinito, ainda que seus membros possam ser finitos, que temos de encontrar uma maneira de tratar.

O método que se segue de definir multiplicação em geral é devido ao dr. Whitehead. É explicado e tratado em detalhe em *Principia Mathematica*, vol.I nota 80 ss, e vol.II, nota 114.

Suponhamos, para começar, que κ é uma classe de classes entre as quais não há nenhuma superposição — digamos os distritos eleitorais num país em que não há voto plural, cada distrito sendo considerado uma classe de eleitores. Procuremos agora escolher um termo em cada classe para ser seu *representante*, como distritos eleitorais fazem quando elegem

membros do Parlamento, supondo que por lei cada distrito tenha de eleger um homem que seja um eleitor naquele distrito. Chegamos assim a uma classe de representantes, que compõem nosso Parlamento, cada um selecionado em um distrito eleitoral. Quantos Paramentos diferentes podem ser escolhidos? Cada distrito eleitoral pode escolher qualquer um de seus eleitores, e, portanto, se houver μ eleitores num distrito, poderá fazer μ escolhas. As escolhas dos diferentes distritos são independentes; assim é óbvio que, quando o número total de distritos for finito, o número de Paramentos possíveis é obtido multiplicando-se juntos os números de eleitores nos vários distritos eleitorais. Quando não soubermos se o número de distritos é finito ou infinito, podemos considerar que o número de Paramentos possíveis *define* o produto dos números dos diferentes distritos eleitorais. Esse é o método pelo qual produtos infinitos são definidos. Podemos agora abandonar nossa ilustração e passar a afirmações exatas.

Suponhamos que κ é uma classe de classes, e suponhamos para começar que não há nenhuma superposição entre duas classes, isto é, que se α e β foram dois diferentes membros de κ , nenhum membro de uma será membro da outra. Chamaremos uma classe uma “seleção” a partir de κ quando ela consistir em apenas um termo de cada membro de κ ; isto é, μ é uma seleção a partir de κ se cada membro de μ pertencer a um membro de κ , e se α for membro de κ , μ e α terão exatamente um termo em comum. Chamaremos a classe de todas as “seleções” a partir de κ a “classe multiplicativa” de κ . O número de termos na classe multiplicativa de κ , isto é, o número de seleções possíveis a partir de κ , será definido como o produto dos números dos membros de κ . Essa definição é igualmente aplicável quer κ seja finito ou infinito.

Antes que possamos ficar inteiramente satisfeitos com essas definições, devemos remover a restrição de que não deve haver nenhuma superposição entre dois termos de κ . Para esse fim, em vez de definir primeiro uma classe chamada uma “seleção”, definiremos primeiro uma relação que chamaremos um “seletor”. Uma relação R será chamada um “seletor” proveniente de κ se, de cada membro de κ , ele escolher um termo como representante daquele membro, isto é, se, dado qualquer membro α de κ , houver apenas um termo x que seja membro de α e tenha a relação R com α ; e R fará unicamente isso. A definição formal é: um “seletor” proveniente de uma classe de classes κ é uma relação um-muitos, tendo κ

por seu domínio inverso, e tal que, se x tiver essa relação com α , então x será membro de α .

Se R for um seletor proveniente de κ , α for membro de κ , e x for o termo que tem a relação R com α , chamaremos x o “representante” de α com respeito à relação R .

Uma “seleção” a partir de κ será agora definida como o domínio de um seletor; e a classe multiplicativa, como antes, será a classe das seleções.

Quando os membros de κ se superpõem, pode haver mais seletores que seleções, visto que um termo x que pertence a duas classes α e β pode ser selecionado uma vez para representar α e uma vez para representar β , dando origem a dois diferentes seletores nos dois casos, mas à mesma seleção. Para fins de definir a multiplicação, é de seletores, não de seleções, que precisamos. Assim definimos: “o produto do número dos membros de uma classe de classes κ ” é o número de seletores provenientes de κ .

Podemos definir exponenciação mediante uma adaptação do plano descrito. Poderíamos, é claro, definir μ^v como o número de seletores provenientes de v classes, cada um dos quais tem μ termos. Mas há objeções a essa definição, porque, se ela for adotada, o axioma multiplicativo (do qual falaremos em breve) será desnecessariamente envolvido. Em vez disso, adotamos a seguinte construção: suponhamos que α é uma classe com μ termos, e β uma classe com v termos.

Suponhamos que y seja membro de β e forme a classe de todos os pares bem ordenados que têm y por seu segundo termo e um membro de α por seu primeiro termo. Haverá μ desses pares para um dado y , uma vez que qualquer membro de α pode ser escolhido para o primeiro termo, e α tem μ membros. Se formarmos agora todas as classes desse tipo que resultam da variação de y , obteremos ao todo v classes, pois y pode ser qualquer membro de β , e β tem v membros. Essas classes v são, cada uma delas, uma classe de pares, a saber, todos os pares que podem ser formados com um membro variável de α e um membro fixo de β . Definimos μ^v como o número de seletores provenientes da classe que consiste nessas v classes. Poderíamos igualmente definir μ^v como o número de seleções, já que, como nossas classes de pares são mutuamente exclusivas, o número de seletores é igual ao número de seleções. Uma seleção a partir de nossa classe de classes será um conjunto de pares ordenados, dos quais haverá exatamente um tendo qualquer membro dado de β por seu segundo termo,

e o primeiro termo pode ser qualquer membro de α . Assim μ^v é definido pelos seletores provenientes de um certo conjunto de v classes, cada uma tendo μ termos, mas o conjunto tem certa estrutura e uma composição mais manejável do que em geral é o caso. A relevância disto para o axioma multiplicativo aparecerá em breve.

O que se aplica à exponenciação aplica-se também ao produto de dois cardinais. Poderíamos definir “ $\mu \times v$ ” como a soma dos números de v classes, cada uma tendo μ termos, mas preferimos defini-lo como o número de pares ordenados a serem formados consistindo em um membro de α seguido por um membro de β , onde α tem μ termos e β tem v termos. Essa definição tem também o objetivo de escapar à necessidade de admitir o axioma multiplicativo.

Com nossas definições, podemos provar as leis formais usuais da multiplicação e exponenciação, mas não podemos provar que um produto só é zero quando um dos fatores é zero. É possível provar isso quando o número de fatores é finito, mas não quando é infinito. Em outras palavras, não podemos provar que, dada uma classe de classes nenhuma das quais é nula, deve haver seletores provenientes delas; ou que, dada uma classe de classes mutuamente exclusivas, deve haver pelo menos uma classe consistindo em um termo proveniente de cada uma das classes dadas. Essas coisas não podem ser provadas; e embora à primeira vista pareçam obviamente verdadeiras, a reflexão gera gradualmente dúvida crescente, até que nos contentamos em registrar a suposição e suas conseqüências, como registramos o axioma das paralelas, sem supor que podemos saber se é verdadeiro ou falso. A suposição, frouxamente formulada, é que seletores e seleções existem quando devemos esperar que existam. Há muitas maneiras equivalentes de afirmar isso com precisão. Podemos começar com a seguinte: “Dada qualquer classe de classes mutuamente exclusivas, das quais nenhuma é nula, há pelo menos uma classe que tem exatamente um termo em comum com cada uma das classes dadas.”

Podemos chamar essa proposição o “axioma multiplicativo”.¹ Primeiro daremos várias formas equivalentes da proposição, e depois consideraremos certas maneiras pelas quais sua verdade ou falsidade é de interesse para matemáticos.

O axioma multiplicativo é equivalente à proposição de que um produto só é zero quando pelo menos um de seus fatores é zero; isto é, que, se

qualquer número de números cardinais for multiplicado juntos, o resultado não pode ser 0, a menos que um dos números envolvidos seja 0.

O axioma multiplicativo é equivalente à proposição de que, se R for qualquer relação, e κ qualquer classe contida no domínio inverso de R , haverá pelo menos uma relação um-muitos que implique R e que tenha κ por seu domínio inverso.

O axioma multiplicativo é equivalente à suposição de que, se α for qualquer classe, e κ todas as subclasses de α com exceção da classe nula, haverá ao menos um seletor proveniente de κ . Foi sob essa forma que o axioma foi primeiro levado à atenção do mundo culto por Zermelo, em seu “Beweis, dass jede Menge wohlgeordnet werden”.² Zermelo vê o axioma como uma verdade inquestionável. Deve ser confessado que, até que ele o explicitasse, os matemáticos o haviam usado sem nenhum escrúpulo, embora se tenha a impressão de que o haviam feito inconscientemente. O crédito devido a Zermelo por tê-lo tornado explícito é inteiramente independente da questão de ser ele verdadeiro ou falso.

Na prova mencionada, Zermelo mostrou que o axioma multiplicativo é equivalente à proposição de que toda classe pode ser bem ordenada, isto é, pode ser arranjada numa série em que toda subclasse tem um primeiro termo (exceto, é claro, a classe nula). A prova completa dessa proposição é difícil, mas não é difícil ver o princípio geral do qual resulta. Ela usa a forma que chamamos “axioma de Zermelo”, isto é, supõe que, dada qualquer classe α , há pelo menos uma relação R um-muitos cujo domínio inverso consiste em todas as subclasses existentes de α e que é tal que, se x tiver a relação R com ξ , então x será membro de ξ . Tal relação tira um “representante” de cada subclasse; freqüentemente acontecerá, é claro, que duas subclasses tenham o mesmo representante. O que Zermelo faz, de fato, é contar os membros de α , um a um, por meio de R e indução transfinita. Colocamos primeiro o representante de α ; vamos chamá-lo x_1 . Depois tomamos o representante da classe que consiste em todo α exceto x_1 ; vamos chamá-lo x_2 . Ele deve ser diferente de x_1 , porque todo representante é membro de sua classe, e x_1 está excluído dessa classe. Procedemos da mesma maneira para retirar x_2 , e deixamos x_2 ser o representante do que resta. Dessa maneira obtemos primeiro uma progressão $x_1, x_2, \dots, x_n, \dots$, supondo que α não é finito. Depois retiramos toda a progressão; deixamos x_ω ser o representante do que resta de α .

Dessa maneira podemos prosseguir até que nada reste. Os sucessivos representantes formarão uma série bem ordenada contendo todos os membros de α . (O que apresentamos é, obviamente, apenas uma indicação das linhas gerais da prova.) Essa proposição é chamada “teorema de Zermelo”.

O axioma multiplicativo é também equivalente à suposição de que de qualquer de dois cardinais que não sejam iguais, um deve ser o maior. Se o axioma for falso, haverá cardinais μ e ν tais que μ não é nem menor, nem igual, nem maior do que ν . Vimos que n_1 e 2^n_0 formam possivelmente um caso de par como esse.

Muitas outras formas do axioma poderiam ser dadas, mas as aqui apresentadas são as mais importantes entre as conhecidas atualmente. Quanto à verdade ou falsidade do axioma em qualquer de suas formas, nada se sabe hoje.

As proposições que dependem do axioma, sem serem reconhecidamente equivalentes a ele, são numerosas e importantes. Tomemos primeiro a conexão da adição e multiplicação. É natural pensarmos que a soma de ν classes mutuamente exclusivas, cada uma tendo μ termos, deve ter $\mu \times \nu$ termos. Quando ν é finito, isso pode ser provado. Mas quando ν é infinito, não pode ser provado sem o axioma multiplicativo, exceto onde, em virtude de alguma circunstância especial, a existência de certos seletores pode ser provada. O axioma multiplicativo é usado da seguinte maneira: suponhamos que temos dois conjuntos de μ classes mutuamente exclusivas, cada um com μ termos, e desejamos provar que a soma de um conjunto tem tantos termos quanto a soma do outro. Para provar isso, devemos estabelecer uma relação um-um. Ora, como há ν classes em cada caso, há alguma relação um-um entre os dois conjuntos de classes; mas o que queremos é uma relação um-um entre seus termos. Consideremos alguma relação um-um S entre as classes. Nesse caso, se κ e λ forem os dois conjuntos de classes, e α for membro de κ , haverá um membro β de λ que será o correlato de α com respeito a S . Ora α e β têm termos μ e, portanto, são similares. Conseqüentemente, há correlações um-um de α e β . O problema é haverem tantas. Para obter uma correlação um-um da soma de κ com a soma de λ , temos de retirar *uma seleção* de um conjunto de classes de correlatores, uma classe do conjunto sendo todos os correlatores um-um de α com β . Se κ e λ forem finitos, em geral não poderemos saber se semelhante seleção existe, a menos que possamos

saber que o axioma multiplicativo é verdadeiro. Portanto, não podemos estabelecer o tipo usual de conexão entre adição e multiplicação.

Esse fato tem várias conseqüências curiosas. Para começar, sabemos que $n_0^2 = n_0 \times n_0 = n_0$. Comumente infere-se disso que a soma de n_0 classes, cada qual com n_0 membros, deve ter ela própria n_0 membros, mas esta inferência é falaciosa, visto que não sabemos se o número de termos em tal soma é $n_0 \times n_0$, nem conseqüentemente se é n_0 . Isso tem uma relação com a teoria dos ordinais transfinitos. É fácil provar que um ordinal que tem n_0 predecessores deve ser um ordinal do que Cantor chama a “segunda classe”, isto é, tal que uma série que tenha esse número ordinal terá n_0 termos em seu campo. É também fácil ver que, se tomamos qualquer progressão de ordinais da segunda classe, os predecessores do limite deles formam no máximo a soma de n_0 classes, cada uma tendo n_0 termos. Infere-se, portanto — falaciosamente, a menos que o axioma multiplicativo seja verdadeiro —, que os predecessores do limite são n_0 em número, e por conseguinte que o limite é um número da “segunda classe”. Isto é, fica supostamente provado que qualquer progressão de ordinais da segunda classe tem um limite que é por sua vez um ordinal da segunda classe. Essa proposição, com o corolário de que ω_1 (o menor ordinal da terceira classe) não é o limite de nenhuma progressão, está envolvida na maior parte da teoria dos ordinais da segunda classe reconhecida. Diante do modo como o axioma multiplicativo é envolvido, a proposição e seu corolário não podem ser considerados provados. Podem ser verdadeiros, ou não. Tudo o que se pode dizer presentemente é que não sabemos. Assim, a maior parte da teoria dos ordinais da segunda classe deve ser considerada não provada.

Uma outra ilustração pode ajudar a elucidar a questão. Sabemos que $2 \times n_0 = n_0$. Poderíamos portanto, supor que a soma de n_0 pares deve ter n_0 termos. Mas, embora possamos provar que isso ocorre às vezes, não podemos provar que ocorre *sempre*, a menos que admitamos o axioma multiplicativo. Isso é ilustrado pelo milionário que comprava um par de meias sempre que comprava um par de botas, e nunca em qualquer outra ocasião, e que tinha tal paixão por comprar que terminou por ter n_0 pares de botas e n_0 pares de meias. O problema é: quantas botas ele tinha, e quantas meias? Naturalmente suporíamos que tinha duas vezes mais botas e duas vezes mais meias que o número de pares que possuía de uma coisa

e de outra, e que portanto tinha n_0 de ambas, já que esse número não aumenta quando é duplicado. Mas esse é um caso da dificuldade, já observada, de conectar a soma de v classes, cada uma tendo μ termos, com $\mu \times v$. Às vezes isso pode ser feito, às vezes não. Em nosso caso, pode ser feito com as botas, mas não com as meias, a não ser por meio de algum estratagema muito artificial. A razão para a diferença é: entre as botas, podemos distinguir direita e esquerda, e portanto podemos fazer uma seleção de uma a partir de cada par, a saber, podemos escolher todas as botas direitas ou todas as botas esquerdas; com as meias, porém, nenhum princípio de seleção como esse se apresenta, e não podemos ter certeza, a menos que admitamos o axioma multiplicativo, de que há alguma classe composta de uma meia retirada de cada par. Daí o problema.

Podemos formular a questão de uma outra maneira. Para provar que uma classe tem n_0 termos, é necessário e suficiente encontrar alguma maneira de arranjar seus termos numa progressão. Não há dificuldade em fazer isso com as botas. Os *pares* são dados como formando um n_0 e, portanto, como o campo de uma progressão. Em cada par, tome a bota esquerda em primeiro lugar e a bota direita em segundo, mantendo a ordem do par inalterada; dessa maneira obtemos uma progressão de todas as botas. No caso das meias, porém, teremos de escolher arbitrariamente, em cada par, qual pé pôr em primeiro lugar; e um número infinito de escolhas arbitrárias é uma impossibilidade. A menos que possamos encontrar uma *regra* para selecionar, isto é, uma relação que seja um seletor, não sabemos se uma seleção é sequer teoricamente possível. É claro que, no caso de objetos no espaço, como meias, podemos sempre encontrar algum princípio de seleção. Por exemplo, podemos tomar os centros de massa das meias: haverá pontos p no espaço tais que, com qualquer par, os centros de massa das duas meias não estarão ambos a igual distância de p ; assim poderemos escolher, de cada par, aquela meia cujo centro de massa está mais próxima de p . Mas não há nenhuma razão teórica para que um método de seleção como esse deva ser sempre possível, e o caso das meias, com um pouco de boa vontade da parte do leitor, pode servir para mostrar como uma seleção poderia ser impossível.

Convém observar que, se *fosse* impossível selecionar um pé de meia de cada par, disso se seguiria que meias não *poderiam* ser arranjadas numa progressão, e portanto que não haveria n_0 delas. Esse caso ilustra que, se μ é um número infinito, um conjunto de μ pares pode não conter o mesmo

número de termos que um outro conjunto de μ pares; pois, dados n_0 pares de botas, há certamente n_0 botas, mas não podemos ter certeza disso no caso das meias, a menos que admitamos o axioma multiplicativo ou recorramos a algum método geométrico de seleção como o exposto anteriormente.

Outro importante problema envolvendo o axioma multiplicativo é a relação entre reflexividade e não-indutividade. Lembremos que no Capítulo 8 ressaltamos que um número reflexivo deve ser não-indutivo, mas o inverso (pelo que se sabe até o presente) só pode ser provado se admitirmos o axioma multiplicativo. Isso acontece da seguinte maneira:

É fácil provar que uma classe reflexiva é uma classe que contém subclasses com n_0 termos. (A classe pode, é claro, ter ela própria n_0 termos.) Assim temos de provar, se pudermos, que, dada qualquer classe não-indutiva, é possível escolher uma progressão entre seus termos. Ora, não há nenhuma dificuldade em mostrar que uma classe não-indutiva deve conter mais termos que qualquer classe indutiva, ou, o que dá no mesmo, que se α for uma classe não-indutiva e ν for qualquer número indutivo, haverá subclasses de α que terão ν termos. Portanto, podemos formar conjuntos de subclasses finitas de α : primeiro uma classe que não tem nenhum termo; depois classes que têm 1 termo (tantos quantos forem os membros de α), depois classes que têm 2 termos, e assim por diante. Obtemos dessa maneira uma progressão de conjuntos de subclasses, cada conjunto consistindo em todos aqueles que têm certo número finito dado de termos. Até agora não usamos o axioma da multiplicação, apenas provamos que o número de coleções de subclasses de α é um número reflexivo, isto é, que, se μ for o número de membros de α , de modo que 2^μ seja o número de subclasses de α e 2^{2^μ} seja o número de coleções de subclasses, então, contanto que μ seja não indutivo, 2^{2^μ} deve ser reflexivo. Mas afastamo-nos muito do que tínhamos intenção de provar.

Para avançar além desse ponto, devemos empregar o axioma multiplicativo. De cada conjunto de subclasses, escolhamos uma, omitindo as subclasses que consistem apenas na classe nula. Isto é, escolhamos uma subclasse que contenha um termo, α_1 ; uma que contenha dois termos, α_2 ; uma que contenha três, digamos α_3 ; e assim por diante. (Podemos fazer isto se o axioma multiplicativo for admitido; do contrário, não sabemos se podemos sempre fazer isso.) Temos agora uma progressão, $\alpha_1, \alpha_2, \alpha_3, \dots$

de subclasses de α , em vez de uma progressão de coleções de subclasses; portanto, estamos um passo mais próximos de nossa meta. Sabemos agora que, admitindo o axioma multiplicativo, se μ for um número não-indutivo, 2^μ deve ser um número reflexivo.

O próximo passo é observar que, embora não possamos ter certeza de que novos membros de α sejam introduzidos em qualquer estágio especificado na progressão $\alpha_1, \alpha_2, \alpha_3, \dots$ podemos ter certeza de que novos membros continuam sendo introduzidos de tempo em tempo. Ilustremos isso. A classe α_1 , que consiste em um termo, é um novo começo; suponhamos que esse único termo é x_1 ; a classe α_2 , que consiste em dois termos, pode conter ou não x_1 ; se contiver, ela introduz um termo novo; e se não contiver, deve introduzir dois termos novos, digamos x_2, x_3 . Nesse caso, é possível que α_3 consista em x_1, x_2, x_3 , e assim não introduza nenhum termo novo, mas nesse caso α_4 deve introduzir um novo termo. As primeiras v classes $\alpha_1, \alpha_2, \alpha_3, \dots, \alpha_v$ contêm, no máximo, $1 + 2 + 3 + \dots + v$ termos, isto é, $v(v+1)/2$ termos; seria portanto possível, se não houvesse repetições nas primeiras v classes, prosseguir com repetições apenas $(v+1) - \text{ésima}$ classe até $[v(v+1)/2] - \text{ésima}$ classe. Nessa altura, porém, os velhos termos não seriam mais suficientemente numerosos para formar uma classe seguinte com o número correto de membros, isto é, $v(v+1)/2+1$, portanto novos termos deveriam ser introduzidos chegar nesse ponto, se não mais cedo. Segue-se que, se omitirmos de nossa progressão $\alpha_1, \alpha_2, \alpha_3, \dots$ todas aquelas classes compostas inteiramente de membros que ocorreram em classes anteriores, ainda teremos uma progressão. Chamemos nossa nova progressão $\beta_1, \beta_2, \beta_3, \dots$ (Teremos $\alpha_1 = \beta_1$ e $\alpha_2 = \beta_2$, porque α_1 e α_2 têm de introduzir novos termos. Podemos ter ou não $\alpha_3 = \beta_3$, mas, em geral, β_μ será α_v , onde v é algum número maior do que μ ; isto é, os β 's são *alguns* dos α 's.) Ora esses β 's são tais que qualquer deles, digamos β_μ , contém membros que não ocorreram em nenhum dos β 's anteriores. Chamemos a parte de β_μ que consiste em novos membros γ . Assim obtemos uma nova progressão $\gamma_1, \gamma_2, \gamma_3, \dots$ (Novamente, γ_1 será idêntico a β_1 e a α_1 ; se α_2 não contiver o único membro de α_1 , teremos $\gamma_2 = \beta_2 = \alpha_2$, mas se α_2 contiver esse único membro, γ_2 consistirá no outro membro de α_2 .) Essa nova progressão de γ 's consiste em classes mutuamente exclusivas. Uma seleção a partir delas será, portanto, uma progressão; isto

é, se x_1 for membro de γ_1 , x_2 será membro de γ_2 , x_3 será membro de γ_3 , e assim por diante; portanto, x_1, x_2, x_3 é uma progressão e é uma subclasse de α . Admitindo o axioma multiplicativo, tal seleção pode ser feita. Assim, usando duas vezes esse axioma podemos provar que, se o axioma for verdadeiro, todo cardinal não-indutivo deve ser reflexivo. Isso poderia ser deduzido também do teorema de Zermelo, segundo o qual, se o axioma for verdadeiro, toda classe pode ser bem ordenada; pois um série bem ordenada deve ter em seu campo um número de termos ou finito ou reflexivo.

O raciocínio direto anterior tem uma vantagem sobre dedução a partir do teorema de Zermelo por não exigir a verdade universal do axioma multiplicativo, mas somente sua verdade tal como aplicado a um conjunto de n_0 classes. Pode acontecer que o axioma se sustente para n_0 classes, mas não para números maiores de classes. Por essa razão é melhor, quando possível, nos contentarmos com a admissão mais restrita. A admissão feita no raciocínio direto anterior é que um produto de n_0 fatores nunca é zero, a menos que um dos fatores seja zero. Podemos formular essa suposição na forma “ n_0 é um número *multiplicável*”, onde um número v é definido como “multiplicável” quando um produto de v fatores nunca for zero, a menos que um dos fatores seja zero. Podemos *provar* que um número *finito* é sempre multiplicável, mas não podemos provar o mesmo com relação a qualquer número infinito. O axioma multiplicativo é equivalente à suposição de que *todos* os números cardinais são multiplicáveis. Mas para identificar o reflexivo com o não-indutivo, ou para tratar do problema das botas e das meias, ou para mostrar que qualquer progressão de números da segunda classe é da segunda classe, precisamos apenas da suposição muito menor de que n_0 é multiplicável.

Não é improvável que haja muito a ser descoberto com relação aos tópicos discutidos neste capítulo. É possível que se encontrem casos em que proposições que parecem envolver o axioma multiplicativo podem ser resolvidas sem ele. É concebível que se venha a demonstrar que o axioma multiplicativo, em sua forma geral, é falso. Desse ponto de vista, o teorema de Zermelo oferece a melhor perspectiva: *poderia* se provar que o contínuo de alguma série ainda mais densa não pode ter seus termos bem ordenados, o que provaria a falsidade do axioma multiplicativo em virtude do teorema de Zermelo. Até agora, porém, não se descobriu nenhum

método para obter esses resultados e o assunto permanece envolto em obscuridade.

Capítulo 13

O axioma da infinidade e os tipos lógicos

O axioma da infinidade é uma suposição que pode ser enunciada da seguinte maneira: “Se n for qualquer número cardinal indutivo, haverá pelo menos uma classe de indivíduos que tenha n termos.”

Se isso for verdade, segue-se, é claro, que há muitas classes de indivíduos que têm n termos, e que o número total dos indivíduos no mundo não é um número indutivo. Pois, pelo axioma, há pelo menos uma classe contendo $n + 1$ termos, do que se segue que há muitas classes de n termos e que n não é o número total de indivíduos no mundo. Como n é *qualquer* número indutivo, segue-se que o número de indivíduos no mundo deve (se nosso axioma for verdadeiro) exceder qualquer número indutivo. Tendo em vista o que vimos no capítulo anterior sobre a possibilidade de cardinais que não são nem indutivos nem reflexivos, não podemos inferir de nosso axioma que haja pelo menos n_0 indivíduos, a menos que admitamos o axioma multiplicativo. Mas não sabemos que há pelo menos n_0 classes de classes, visto que os cardinais indutivos são classes de classes, e formam uma progressão se nosso axioma for verdadeiro. A maneira pela qual a necessidade desse axioma surge pode ser explicada assim: uma das suposições de Peano é que dois cardinais indutivos jamais têm o mesmo sucessor, isto é, não teremos $m + 1 = n + 1$ a menos que $m = n$, se m e n forem cardinais indutivos. No Capítulo 8 tivemos ocasião de usar uma suposição praticamente igual a essa de Peano, a saber, que, se n for um cardinal indutivo, n não será igual a $n + 1$. Poderíamos pensar que é possível provar isso. É possível provar que, se α for uma classe indutiva, e n for o número dos membros de α , então n não será igual a $n + 1$. Essa proposição é facilmente provada por indução, e poderíamos pensar que implica a outra. Mas de fato não o faz, pois poderia não haver tal classe α . O que ela implica é isto: se n for um cardinal indutivo tal que haja pelo

menos uma classe com n membros, então n não é igual a $n + 1$. O axioma da infinidade nos assegura (quer verdadeira ou falsamente) que há classes com n membros, e portanto nos permite afirmar que n não é igual a $n + 1$. Sem esse axioma, porém, ficaríamos com a possibilidade de n e $n + 1$ serem ambos a classe nula.

Ilustremos essa possibilidade com um exemplo: suponhamos que houvesse exatamente nove indivíduos no mundo. (Quanto ao que se deve entender pela palavra “indivíduo”, peço ao leitor que seja paciente.) Assim, os cardinais indutivos de 0 a 9 seriam tais como esperamos, mas 10 (definido como $9 + 1$) seria a classe nula. Convém lembrar que $n+1$ pode ser definido da seguinte maneira: $n + 1$ é a coleção de todas aquelas classes que têm um termo x tal que, quando x é retirado, resta uma classe de n termos. Aplicando agora essa definição, vemos que, no caso suposto, $9 + 1$ é uma classe que consiste em nenhuma classe, isto é, é a classe nula. O mesmo será verdade acerca de $9 + 2$, ou em geral de $9 + n$, a menos que n seja zero. Assim 10 e todos os cardinais indutivos subsequentes serão todos idênticos, visto que serão todos a classe nula. Num caso como esse, os cardinais indutivos não formarão uma progressão, nem será verdade que dois deles não podem ter o mesmo sucessor, pois 9 e 10 serão ambos sucedidos pela classe nula (10 sendo ele mesmo a classe nula). É para evitar catástrofes aritméticas como essa que precisamos do axioma da infinidade.

De fato, enquanto ficamos satisfeitos com a aritmética dos números integrais finitos, e não introduzimos números inteiros infinitos nem classes ou séries infinitas de números inteiros ou razões finitas, é possível obter todos os resultados desejados sem o axioma da infinidade. Isto é, podemos lidar com a adição, a multiplicação e a exponenciação de números inteiros finitos e de razões, mas não podemos lidar com números inteiros infinitos ou com irracionais. Assim, a teoria do transfinito e a teoria dos números reais não nos ajudam. Devemos agora explicar como esses vários resultados são obtidos.

Admitindo que o número de indivíduos no mundo é n , o número de classes de indivíduos será 2^n . Isso em virtude da proposição geral mencionada no Capítulo 8, segundo a qual o número de classes contido numa classe que tem n membros é 2^n . Ora 2^n é sempre maior do que n . Portanto, o número de classes no mundo é maior que o número de indivíduos. Se supusermos agora que o número de indivíduos é 9, como

fizemos há pouco, o número de classes será 2^9 , isto é, 512. Logo, se tomamos nossos números como aplicados à contagem de classes em vez de à contagem de indivíduos, nossa aritmética será normal até chegarmos a 512: o primeiro número a ser nulo será 513. E se avançarmos para classes de classes, faremos ainda melhor: o número delas será 2^{512} , um número grande a ponto de desconcertar nossa imaginação, pois tem cerca de 153 dígitos. E se avançarmos para classes de classes de classes, obteremos um número representado por 2 elevado a uma potência com cerca de 153 dígitos; o número de dígitos nesse número será cerca de três vezes 10^{152} . Numa época de escassez de papel é indesejável escrever tal número, e se quisermos números maiores podemos ir mais longe na hierarquia lógica. Dessa maneira, podemos fazer com que qualquer cardinal indutivo designado encontre seu lugar entre os números que não são nulos, meramente avançando pela hierarquia por uma distância suficiente.¹

No tocante a razões, temos um estado de coisas muito similar. Para que uma razão μ/ν tenha as propriedades esperadas, deve haver objetos suficientes do tipo, seja ele qual for, que esteja sendo contado, para assegurar que a classe nula não intrometa subitamente. Mas isso pode ser assegurado, para qualquer razão μ/ν , sem o axioma da infinidade, meramente avançando uma distância suficiente pela hierarquia acima. Se não conseguirmos isso contando indivíduos, podemos tentar contando classes de indivíduos; se ainda não conseguirmos, podemos tentar classes de classes, e assim por diante. Em última análise, por menos indivíduos que haja no mundo, alcançaremos um estágio em que haverá muito mais do que μ objetos, seja que número indutivo μ possa ser. Mesmo que não houvesse absolutamente nenhum indivíduo, isso ainda seria verdade, pois haveria então uma classe, a saber, a classe nula, duas classes de classes (a saber, a classe nula de classes e a classe cujo único membro seria a classe nula de indivíduos), quatro classes de classes de classes, 16 no estágio seguinte, 65.536 no seguinte, e assim por diante. Nenhuma suposição como o axioma da infinidade é, portanto, necessária para se alcançar qualquer razão dada ou qualquer cardinal indutivo dado.

É quando desejamos lidar com toda a classe ou série de cardinais indutivos ou de razões que o axioma é necessário. Precisamos de toda a classe de cardinais indutivos para estabelecer a existência de n_0 , e de toda a série para estabelecer a existência de progressões: para esses resultados,

precisamos ser capazes de fazer uma única classe ou série em que nenhum cardinal indutivo seja nulo. Precisamos de toda a série de razões em ordem de magnitude para definir números reais como segmentos: essa definição não dará o resultado desejado a menos que a série de razões seja compacta, o que ela não pode ser se o número total de razões, no estágio envolvido, for finito.

Seria natural supor — como eu próprio supus outrora — que, por meio de construções tais como as que temos considerado, o axioma da infinidade poderia ser *provado*. Pode-se dizer: suponhamos que o número de indivíduos é n , onde n pode ser 0 sem estragar nosso raciocínio; então se formarmos o conjunto completo de indivíduos, classes, classes de classes etc., todos tomados juntos, o número de termos em todo o nosso conjunto será

$$n + 2^n + 2^{2^n} \dots \text{ao infinito},$$

que é n_0 . Assim, tomando todos os tipos de objetos juntos, e não nos limitando a objetos de nenhum tipo, obteremos certamente uma classe infinita, e portanto não precisaremos do axioma da infinidade. Isso é o que se poderia dizer.

Mas, antes de analisar esse raciocínio, a primeira coisa a observar é que ele dá certa impressão de prestidigitação: alguma coisa nos faz lembrar o mágico que tira coisas do chapéu. O homem que emprestou o chapéu tem absoluta certeza de que não havia um coelho vivo nele antes, e não tem a menor idéia de como o coelho foi parar lá. Assim também o leitor, se tiver um senso de realidade robusto, se sentirá convencido de que é impossível produzir uma coleção infinita a partir de uma coleção finita de indivíduos, mesmo que seja incapaz de dizer onde está a falha da construção apresentada. Seria um erro dar muita ênfase a essas impressões de prestidigitação; como outras emoções, elas podem facilmente nos enganar. Mas elas oferecem uma base *prima facie* para se examinar com muita atenção qualquer raciocínio que as suscite. E quando o raciocínio anterior for examinado, ele se revelará, na minha opinião, falacioso, embora se trate de uma falácia sutil e de modo algum fácil de evitar coerentemente.

A falácia envolvida é a que pode ser chamada “confusão de tipos”. Para explicar o assunto dos “tipos” por completo seria necessário um volume inteiro; além disso, o objetivo deste livro é evitar aquelas partes dos

assuntos que ainda continuam obscuras e controversas, isolando, para a conveniência de iniciantes, as que podem ser aceitas como corporificando verdades matematicamente verificadas. Ora, a teoria dos tipos não pertence, enfaticamente, à parte acabada e certa de nosso assunto: grande parte dessa teoria continua incompleta, confusa e obscura. Mas a necessidade de *alguma* doutrina de tipo é menos duvidosa que a forma precisa que ela deveria assumir; e em conexão com o axioma da infinidade é particularmente fácil ver a necessidade de alguma doutrina.

Essa necessidade resulta, por exemplo, da “contradição do maior cardinal”. Vimos no Capítulo 8 que o número de classes contido numa dada classe é sempre maior do que o número de membros da classe, e inferimos que não há o maior número cardinal. Mas se pudéssemos, como sugerimos um momento atrás, somar numa classe os indivíduos, as classes de indivíduos, as classes de classes de indivíduos etc., obteríamos uma classe da qual suas próprias subclasses seriam membros. A classe que consistiria em todos os objetos que podem ser contados, de qualquer tipo, deve, para poder existir, ter um número cardinal que seja o maior possível. Como todas as suas subclasses serão membros dela, o número de subclasses não pode ser maior do que o número de membros. Chegamos, portanto, a uma contradição.

Quando cheguei pela primeira vez a essa contradição, em 1901, tentei descobrir alguma falha na prova de Cantor de que não há o maior número cardinal, que demos no Capítulo 8. Aplicando essa prova à suposta classe de todos os objetos imagináveis, fui levado a uma nova e mais simples contradição.

A classe completa que estamos considerando, que deve abarcar todas as coisas, deve abarcar a si mesma como um de seus membros. Em outras palavras, se houver algo como “todas as coisas”, então “todas as coisas” é alguma coisa e é membro da classe de “todas as coisas”. Mas normalmente uma classe não é membro de si mesma. A humanidade, por exemplo, não é um homem. Formemos agora a reunião de todas as classes que não são membros de si mesmas. Esta é uma classe: ela é membro de si mesma ou não? Se for, é uma das classes que não são membros de si mesmas, isto é, não é membro de si mesma. Se não, não é uma das classes que não são membros de si mesma, isto é, é membro de si mesma. Portanto as duas hipóteses — isto é, que ela é e que não é membro de si mesma — implicam ambas sua contraditória. Isso é uma contradição.

Não há dificuldade em produzir contradições similares *ad libitum*. A solução de tais contradições pela teoria dos tipos é formulada em *Principia Mathematica*,² e também, mais brevemente, em artigos do presente autor no *American Journal of Mathematics*³ e na *Revue de Metaphysique et de Morale*.⁴ No momento, um esboço da solução deve bastar.

A falácia consiste na formação do que podemos chamar classes “impuras”, isto é, classes que não são puras em relação a “tipo”. Como veremos num capítulo posterior, classes são ficções lógicas, e uma afirmação que parece ser sobre uma classe só será significativa se for passível de tradução numa forma em que não se faça nenhuma menção à classe. Isso impõe uma limitação aos modos como o que são nominalmente, embora não na realidade, nomes para classes podem ocorrer significativamente: uma sentença ou conjunto de símbolos em que tais pseudonomes ocorrem de maneiras erradas não é falsa, mas estritamente desprovida de sentido. A suposição de que uma classe é, ou de que não é, membro de si mesma é desprovida de sentido exatamente dessa maneira. E, de maneira mais geral, supor que uma classe de indivíduos é membro, ou não é membro, de outra classe de indivíduos será supor um absurdo; e construir simbolicamente qualquer classe cujos membros não sejam todos do mesmo grau na hierarquia lógica é usar símbolos de uma maneira que os torna não mais simbólicos de coisa alguma.

Assim se houver n indivíduos no mundo, e 2^n classes de indivíduos, não podemos formar uma nova classe, consistindo tanto em indivíduos quanto em classes e tendo $n + 2^n$ membros. Dessa maneira, a tentativa de escapar da necessidade do axioma da infinidade se frustra. Não pretendo ter explicado a doutrina dos tipos, ou ter feito mais que indicar, grosseiramente, por que há necessidade de tal doutrina. Meu objetivo foi apenas dizer o necessário para mostrar que não podemos *provar* a existência de números e classes infinitos por esses métodos de mágico que estivemos examinando. Restam, contudo, outros métodos possíveis, que devem ser considerados.

Vários raciocínios que professam provar a existência de classes infinitas são dados em *Principles of Mathematics*, § 339 (p.357). Na medida em que esses raciocínios supõem que, se n for um cardinal indutivo, n não será igual a $n + 1$, já tratamos deles. Há um raciocínio, sugerido por uma

passagem de *Parmênides* de Platão, segundo o qual se há um número como 1, então 1 tem existência; mas 1 não é idêntico à existência e, portanto, 1 e existência são dois, e, portanto, há um número como 2, e dois juntamente com 1 e existência dão uma classe de três membros, e assim por diante. Esse raciocínio é falacioso, em parte porque “existência” não é um termo que tenha qualquer sentido definido, e ainda mais porque, se um sentido definido fosse inventado para ele, verificaríamos que os números não têm existência — eles são, de fato, o que é chamado “ficção lógica”, como veremos quando passarmos a considerar a definição de classe.

O raciocínio de que o número de números de 0 a n (ambos incluídos) é $n + 1$ depende da suposição de que até e incluindo n nenhum número é igual a seu sucessor, o qual, como vimos, nem sempre será verdadeiro se o axioma da infinidade for falso. É preciso compreender que a equação $n = n + 1$, que poderia ser verdadeira para um n finito se n excedesse o número total de indivíduos no mundo, é muito diferente da mesma equação tal como aplicada a um número reflexivo. Tal como aplicada a um número reflexivo, ela significa que, dada uma classe de n termos, essa classe será “similar” àquela obtida adicionando-se um outro termo. Mas tal como aplicada a um número grande demais para o mundo real, ela significará meramente que não há nenhuma classe de n indivíduos, e nenhuma classe de $n + 1$ indivíduos; ela não significa que, se subirmos pela hierarquia de tipos o bastante para assegurar a existência de uma classe de n termos, verificaremos então que essa classe é “similar” àquela de $n + 1$ termos, pois se n for indutivo esse não será o caso, de maneira absolutamente independente da verdade ou falsidade do axioma da infinidade.

Há um raciocínio empregado tanto por Bolzano⁵ quanto por Dedekind⁶ para provar a existência de classes reflexivas. O raciocínio, em suma, é o seguinte: um objeto não é idêntico à idéia do objeto, mas há (pelo menos no reino do ser) uma idéia de cada objeto. A relação de um objeto com a idéia é um-um, e as idéias são somente alguns entre os objetos. Portanto, a relação “idéia de” constitui um reflexo de toda a classe dos objetos numa parte de si mesma, a saber, naquela parte que consiste em idéias. Assim, a classe dos objetos e a classe das idéias são ambas infinitas. Esse raciocínio é interessante, não só em si mesmo, mas porque os erros que contém (ou o que julgo serem erros) são de um tipo instrutivo. O principal erro consiste em supor que há uma idéia de cada objeto. É extremamente difícil, é claro, decidir o que se deve entender por uma “idéia”; mas suponhamos que

sabemos. Devemos então supor que, começando (digamos) com Sócrates, há a idéia de Sócrates, e assim por diante *ad infinitum*. Ora, fica claro que esse não é o caso no sentido de que todas essas idéias têm existência empírica real na mente das pessoas. Além do terceiro ou do quarto estágio, elas se tornam míticas. Para que o raciocínio se mantenha, as “idéias” invocadas devem ser idéias platônicas armazenadas no céu, pois certamente não estão na Terra. Mas nesse caso torna-se imediatamente duvidoso que tais idéias existam. Só poderíamos saber que elas existem com base em alguma teoria lógica que provasse que é necessário que para cada coisa exista uma idéia dessa coisa. Certamente não podemos obter esses resultados empiricamente, ou aplicá-lo, como o faz Dedekind, a “*meine Gedankenwelt*” — o mundo de meus pensamentos.

Se estivéssemos interessados em examinar completamente a relação entre idéia e objeto, teríamos de encetar várias investigações psicológicas e lógicas que não são relevantes para nosso principal objetivo. Mas alguns pontos adicionais precisam ser destacados. Se “idéia” deve ser compreendida logicamente, ela pode ser *idêntica* ao objeto, ou pode representar uma *descrição* (no sentido a ser explicado num capítulo posterior). No primeiro caso, o raciocínio malogra, pois era essencial para a prova da reflexividade que objeto e idéia fossem distintos. No segundo caso, o raciocínio também malogra, porque a relação entre objeto e descrição não é um-um: há inúmeras descrições corretas para qualquer objeto dado. Por exemplo, Sócrates pode ser descrito como “o mestre de Platão”, ou como “o filósofo que tomou cicuta”, ou como “o marido de Xantipa”. Se — para tomar a hipótese que resta — devemos interpretar “idéia” psicologicamente, é preciso afirmar que não há nenhuma entidade psicológica definida que poderia ser chamada *a* idéia do objeto: há inúmeras crenças e atitudes, cada uma das quais poderia ser chamada uma idéia do objeto no sentido em que poderíamos dizer “minha idéia de Sócrates é muito diferente da sua”, mas não há nenhuma entidade central (exceto o próprio Sócrates) para ligar entre si várias “idéias de Sócrates” e, portanto, não há essa relação um-um entre idéia e objeto que o raciocínio supõe. Como já observamos, não é tampouco psicologicamente verdadeiro, é claro, que haja idéias (por mais extenso que seja o sentido do termo) de mais do que uma proporção minúscula das coisas no mundo. Por todas essas razões, esse raciocínio a favor da existência lógica de classes reflexivas deve ser rejeitado.

Poder-se-ia pensar que, o que quer que possa ser dito sobre raciocínios *lógicos*, os raciocínios *empíricos* deriváveis do espaço e tempo, da diversidade das cores etc. são inteiramente suficientes para provar a existência real de um número infinito de particulares. Não acredito nisso. Não tenho nenhuma razão, afora uma idéia preconcebida, para acreditar na extensão infinita do espaço e tempo, pelo menos no sentido em que espaço e tempo são fatos físicos, não ficções matemáticas. Naturalmente encaramos espaço e tempo como contínuos, ou, pelo menos, como compactos; mas isso é novamente sobretudo idéia preconcebida. A teoria dos *quanta* na física, quer seja verdadeira ou falsa, ilustra o fato de que a física não pode nunca fornecer prova de continuidade, embora muito possivelmente possa fornecer prova em contrário. Os sentidos não são suficientemente exatos para distinguir entre movimento contínuo e rápida sucessão discreta, como qualquer pessoa pode descobrir num cinema. Um mundo em que todo movimento consistisse numa série de pequenos solavancos finitos seria empiricamente indistinguível de um em que o movimento fosse contínuo. Tomaria espaço demais defender essas teses adequadamente; por enquanto estou meramente sugerindo-as à consideração do leitor. Se forem válidas, segue-se que não há razão empírica para se acreditar que o número de particulares no mundo é infinito, e que possa ser algum dia; também não há, no momento, nenhuma razão empírica para se acreditar que esse número é finito, embora seja teoricamente concebível que algum dia possa haver indícios apontando, ainda que não conclusivamente, nessa direção.

Do fato de que o infinito não é contraditório, mas também não é demonstrável logicamente, devemos concluir que nada pode ser conhecido *a priori* com relação ao caráter finito ou infinito das coisas no mundo. A conclusão é, portanto, para adotar a fraseologia leibniziana, que alguns dos mundos possíveis são finitos, alguns são infinitos, e não temos nenhum meio de saber a qual desses dois tipos nosso mundo real pertence. O axioma da infinidade será verdadeiro em alguns mundos e falso em outros; se é verdadeiro ou falso neste mundo, não temos como saber.

Ao longo deste capítulo os sinônimos “indivíduo” e “particular” foram usados sem explicação. Seria impossível explicá-los adequadamente sem um estudo mais longo da teoria dos tipos do que seria apropriado para a presente obra, mas algumas palavras antes de deixarmos esse tópico

podem fazer algo para diminuir a obscuridade que de outro modo envolveria o significado dessas palavras.

Numa afirmação comum, podemos distinguir um verbo, que expressa um atributo ou relação, do substantivo, que expressa o sujeito do atributo ou os termos da relação. “César viveu” confere um atributo a César; “Bruto matou César” expressa uma relação entre Bruto e César. Usando a palavra “sujeito” num sentido geral, podemos chamar tanto Bruto quanto César de sujeitos desta proposição: o fato de Bruto ser gramaticalmente o sujeito e César o objeto é logicamente irrelevante, pois a mesma ocorrência pode ser expressa pelas palavras “César foi morto por Bruto”, em que César é o sujeito gramatical. Assim, no tipo mais simples de proposição teremos um atributo ou relação válida acerca de, ou entre, um, dois ou mais “sujeitos” no sentido amplo. (Uma relação pode ter mais de dois termos; por exemplo, “A dá B a C” é uma relação de três termos.) Ora, ocorre muitas vezes que, num exame mais atento, verifica-se que os sujeitos aparentes não são realmente sujeitos, sendo passíveis de análise; o único resultado disso, porém, é que novos sujeitos tomam seus lugares. Acontece também que o verbo possa ser gramaticalmente transformado em sujeito; por exemplo, podemos dizer: “Matar é uma relação válida entre Bruto e César.” Mas nesses casos a gramática é enganosa, e numa afirmação direta, seguindo as regras que deveriam guiar a gramática filosófica, Bruto e César aparecerão como sujeitos e matar como o verbo.

Somos assim levados à concepção de termos que, quando ocorrem em proposições, podem ocorrer somente como sujeitos, e nunca de qualquer outra maneira. Isso é parte da antiga definição escolástica de *substância*; mas a persistência ao longo do tempo, que pertencia a essa noção, não faz parte de maneira alguma da noção em que estamos interessados. Definiremos “nomes próprios” como aqueles termos que só podem ocorrer como *sujeitos* em proposições (usando “sujeito” no sentido ampliado explicado há pouco). Além disso, definiremos “indivíduos” ou “particulares” como os objetos que podem ser nomeados por nomes próprios. (Seria melhor defini-los diretamente, em vez de fazê-lo por meio do tipo de símbolos pelos quais são representados; mas para fazer isso teríamos de mergulhar mais profundamente na metafísica do que é desejável aqui.) É possível, é claro, que haja um retrocesso interminável: que tudo aquilo que aparece como um particular revele-se realmente, a um exame mais atento, uma classe ou algum tipo de complexo. Se esse for o

caso, o axioma da infinidade deve, é claro, ser verdadeiro. Mas se não for, deve ser teoricamente possível para a análise atingir sujeitos últimos, e são esses que dão o sentido de “particulares” ou “indivíduos”. É ao número destes que o axioma da infinidade supostamente se aplica. Se isso for verdadeiro em relação a eles, será verdadeiro em relação a classes deles, e a classes de classes deles, e assim por diante; de maneira similar, se for falso em relação a eles, será falso de um extremo ao outro dessa hierarquia. É natural, portanto, enunciar o axioma em relação a eles e não em relação a qualquer outro estágio na hierarquia. Parece não haver, contudo, nenhum método conhecido para se saber se o axioma é verdadeiro ou falso.

Capítulo 14

Incompatibilidade e a teoria da dedução

Já exploramos, é verdade que um pouco às pressas, aquela parte da filosofia da matemática que não exige um exame crítico da idéia de *classe*. No capítulo anterior, contudo, vimo-nos diante de problemas que tornam tal exame imperativo. Antes que possamos empreendê-lo, devemos considerar certas outras partes da filosofia da matemática que foram ignoradas até este ponto. Num tratamento sintético, as partes de que trataremos agora vêm antes: elas são mais fundamentais que qualquer coisa que tenhamos discutido até este momento. Três tópicos nos interessarão antes que cheguemos à teoria das classes, a saber: (1) a teoria da dedução, (2) funções proposicionais, (3) descrições. Dessas, a terceira não é logicamente pressuposta na teoria das classes, mas é um exemplo mais simples do *tipo* de teoria necessário quando se lida com classes. É o primeiro tópico, a teoria da dedução, que nos ocupará nesse capítulo.

A matemática é uma ciência dedutiva: a partir de certas premissas, ela chega, por um processo estrito de dedução, aos vários teoremas que a constituem. É verdade que, no passado, deduções matemáticas muitas vezes careciam gravemente de rigor; é também verdade que rigor perfeito é um ideal quase inatingível. No entanto, na medida em que falta rigor a uma prova matemática, ela é defeituosa; é inútil insistir em que o senso comum mostra que o resultado está correto, pois se devêssemos confiar nisso, seria melhor prescindir por completo do raciocínio do que recorrer a uma falácia para salvar o senso comum. Nenhum apelo ao senso comum, ou à “intuição”, ou a qualquer coisa exceto lógica dedutiva estrita deveria ser necessário em matemática depois que as premissas foram formuladas.

Kant, após observar que os geômetras de seu tempo não eram capazes de provar seus teoremas unicamente por meio de raciocínio, exigindo um apelo à figura, inventou uma teoria do raciocínio matemático segundo a qual a inferência nunca é estritamente lógica, exigindo sempre o apoio da

chamada “intuição”. Toda a tendência da matemática moderna, com sua maior busca de rigor, foi contra essa teoria kantiana. As coisas na matemática da época de Kant que não podem ser *provadas*, não podem ser *conhecidas* — por exemplo o axioma das paralelas. O que pode ser conhecido, na matemática e por métodos matemáticos, é o que pode ser deduzido da lógica pura. As demais coisas que devem pertencer ao conhecimento humano devem ser verificadas de outra maneira — empiricamente, por intermédio dos sentidos ou de alguma forma de experiência, mas não *a priori*. Os fundamentos positivos para essa tese podem ser encontrados em *Principia Mathematica*, *passim*; uma defesa controversa dela é dada nos *Principles of Mathematics*. Não podemos aqui fazer mais do que remeter o leitor a essas obras, uma vez que o assunto é demasiado vasto para um tratamento apressado. Enquanto isso, suporemos que toda matemática é dedutiva e prosseguiremos para investigar o que está envolvido na dedução.

Na dedução, temos uma ou mais proposições chamadas *premissas*, a partir das quais inferimos uma proposição chamada *conclusão*. Para nossos objetivos, será conveniente, quando houver originalmente várias premissas, amalgamá-las numa única proposição, de modo a poder falar d’*a* premissa bem como d’*a* conclusão. Assim podemos considerar a dedução como um processo pelo qual passamos do conhecimento de uma certa proposição, a premissa, ao conhecimento de uma certa outra proposição, a conclusão. Mas não conceberemos semelhante processo como dedução lógica a menos que ele seja *correto*, isto é, a menos que haja entre premissa e conclusão uma relação tal que tenhamos direito de acreditar na conclusão se soubermos que a premissa é verdadeira. É especialmente essa relação que tem interesse na teoria lógica da dedução.

Para sermos capazes de inferir validamente a verdade de uma proposição, devemos saber que alguma outra proposição é verdadeira, e que há entre as duas uma relação do tipo “implicação”, isto é, que (como dizemos) a premissa “implica” a conclusão. (Definiremos essa relação em breve.) Ou ainda, podemos saber que certa outra proposição é falsa, e que há uma relação entre as duas chamada “disjunção”, expressa por “*p* ou *q*”,¹ de tal modo que o conhecimento de que uma é falsa nos permite inferir que a outra é verdadeira. Assim também o que desejamos inferir pode ser a falsidade de uma proposição, não sua verdade. Isso pode ser inferido da verdade de uma outra proposição, contanto que saibamos que as duas são

“incompatíveis”, isto é, que se uma é verdadeira, a outra é falsa. Ela pode também ser inferida da falsidade de uma outra proposição, exatamente nas mesmas circunstâncias em que a verdade da outra poderia ter sido inferida da verdade de um, isto é, da falsidade de p podemos inferir a falsidade de q quando q implica p . Todos esses quatro casos são casos de inferência. Quando nossas mentes estão fixadas na inferência, parece natural tomar a “implicação” como a relação fundamental primitiva, visto que essa é relação que deve vigorar entre p e q para que possamos inferir a *verdade* de q da *verdade* de p . Por razões técnicas, porém, essa não é a melhor idéia primitiva a escolher. Antes de passamos a idéias primitivas e definições, consideremos mais extensamente as várias funções das proposições sugeridas pelas relações de proposições mencionadas.

A mais simples dessas funções é a negativa, “não- p ”. Essa é aquela função de p que é verdadeira quando p é falso, e falso quando p é verdadeiro. É conveniente falar da verdade de uma proposição, ou de sua falsidade, como seu “valor de verdade”² isto é, *verdade* é o “valor de verdade” de uma proposição verdadeira, e *falsidade* é o valor de verdade de uma proposição falsa. Assim não- p tem o valor de verdade oposto a p .

Podemos tomar em seguida a *disjunção*, “ p ou q ”. Essa é uma função cujo valor de verdade é verdade quando p é verdadeiro e também quando q é verdadeiro, mas é falsidade quando p e q são ambos falsos.

Em seguida podemos tomar a *conjunção*, “ p e q ”. Essa tem a verdade por seu valor de verdade quando p e q são ambos verdadeiros; de outro modo tem falsidade por seu valor de verdade.

Tomemos em seguida a *incompatibilidade*, isto é, “ p e q não são ambos verdadeiros”. Isso é a negação de uma conjunção; é também a disjunção das negações de p e q , ou seja, é “não- p ou não- q ”. Seu valor de verdade é verdade quando p é falso e igualmente quando q é falso; seu valor de verdade é falsidade quando p e q são ambos verdadeiros.

Por último tomemos a *implicação*, isto é, “ p implica q ” ou “se p , então q ”. Isso deve ser compreendido no sentido mais amplo que nos permitirá inferir a verdade de q se conhecermos a verdade de p . Assim, nós a interpretamos como tendo o sentido: “a menos que p seja falso, q é verdadeiro”, ou “ou p é falso ou q é verdadeiro”. (O fato de “implicar” poder ter outros sentidos não nos interessa; esse é o sentido conveniente para nós.) Isto é, “ p implica q ” deve significar “não- p ou q ”: seu valor de

verdade será verdade se p for falso, igualmente se q for verdadeiro, e será falsidade se p for verdadeiro e q for falso.

Temos, portanto, cinco funções: negação, disjunção, conjunção, incompatibilidade e implicação. Poderíamos ter acrescentado outras, por exemplo, “falsidade conjunta, “não- p e não- q ”, mas as cinco apresentadas serão suficientes. A negação difere das outras quatro por ser uma função de *uma* proposição, ao passo que as outras quatro são funções de *duas*. Mas todas as cinco se assemelham nisto: seu valor de verdade depende unicamente do das proposições que são seus argumentos. Dada a verdade ou falsidade de p , ou de p e q (conforme o caso), podemos saber da verdade ou falsidade da negação, disjunção, conjunção, incompatibilidade ou implicação. Uma função de proposições que tem essa propriedade é chamada “função de verdade”.

Todo o significado de uma função de verdade é esgotado pela afirmação das circunstâncias em que ela é verdadeira ou falsa. “Não- p ”, por exemplo, é simplesmente aquela função de p que é verdadeira quando p é falso, e falsa quando p é verdadeiro: nenhum sentido adicional lhe pode ser atribuído. O mesmo se aplica a “ p ou q ” e ao restante. Segue-se que duas funções de verdade que têm o mesmo valor de verdade para todos os valores do argumento são indistinguíveis. Por exemplo, “ p e q ” é a negação de “não p ou não- q ” e vice-versa; assim, qualquer um dos dois pode ser *definido* como a negação do outro. Não há nenhum significado adicional numa função de verdade além das condições nas quais ela é verdadeira ou falsa.

É claro que as cinco funções de verdade mencionadas não são todas independentes. Podemos definir algumas delas em termos de outras. Não há grande dificuldade em reduzir seu número a duas; as duas escolhidas em *Principia Mathematica* são a negação e a disjunção. A implicação é então definida como “não- p ou q ”; a incompatibilidade como “não- p ou não- q ”; a conjunção como a negação da incompatibilidade. Mas foi mostrado por Sheffer³ que podemos nos contentar com *uma* idéia primitiva de todas as cinco, e por Nicod⁴ que isso nos permite reduzir as proposições primitivas requeridas na teoria da dedução a dois princípios não formais e um formal. Para esse propósito, podemos tomar como nosso único indefinível ou a incompatibilidade ou a falsidade conjunta. Escolheremos a primeira.

Nossa idéia primitiva, agora, é uma certa função de verdade chamada “incompatibilidade”, que denotaremos por p/q . A negação pode ser definida de imediato como a incompatibilidade de uma proposição consigo mesma, *i.e.* “não- p ” é definido como p/p . A disjunção é a incompatibilidade de não- p e não- q , isto é, é $(p/p) (q/q)$. A implicação é a incompatibilidade de p e não q , isto é, $p | (q/q)$. A conjunção é a negação da incompatibilidade, ou seja, é $(p/q) | (p/q)$. Assim todas as nossas quatro outras funções são definidas em termos de incompatibilidade.

É óbvio que não há limite para a produção de funções de verdade, seja introduzindo mais argumentos ou repetindo argumentos. O que nos interessa é a conexão desse assunto com inferência.

Se sabemos que p é verdade e que p implica q , podemos avançar e afirmar q . A inferência envolve sempre, inevitavelmente, algo de psicológico: é um método pelo qual chegamos a novo conhecimento, e o que há nela de não-psicológico é a relação que nos permite inferir corretamente; mas a passagem real da asserção p para a asserção de q é um processo psicológico, e não devemos procurar representá-lo em termos puramente lógicos.

Na prática matemática, quando inferimos, temos sempre alguma expressão que contém proposições variáveis, digamos p e q , que sabemos, em virtude de sua forma, ser verdadeira para todos os valores de p e q ; temos também alguma outra expressão, parte da primeira, que também sabemos ser verdadeira para todos os valores de p e q ; e em virtude dos princípios de inferência, somos capazes de abandonar essa parte de nossa expressão original e afirmar o que resta. Essa explicação um pouco abstrata pode ser elucidada mediante alguns exemplos.

Suponhamos que conhecemos os cinco princípios formais da dedução enumerados em *Principia Mathematica*. (M. Nicod os reduziu a um, mas como essa é uma proposição complicada, começaremos com os cinco.) São as seguintes:

(1) “ p ou p ” implica p — isto é, se ou p é verdadeiro ou p é verdadeiro, então p é verdadeiro.

(2) q implica “ p ou q ” — isto é, a disjunção “ p ou q ” é verdadeira quando uma de suas alternativas for verdadeira.

(3) “ p ou q ” implica “ q ou p ”. Isso não seria necessário se tivéssemos uma notação teoricamente mais perfeita, visto que na concepção de disjunção não há ordem envolvida, de modo que “ p ou q ” e “ q ou p ”

deveriam ser idênticos. Mas como nossos símbolos, e qualquer forma conveniente, introduzem inevitavelmente uma ordem, precisamos de suposições adequadas para mostrar que a ordem é irrelevante.

(4) Se ou p é verdadeiro ou “ q ou r ” é verdadeiro, então ou q é verdadeiro ou “ p ou r ” é verdadeiro. (A obliquidade dessa proposição serve para aumentar seu poder dedutivo.)

(5) Se q implica r , então “ p ou q ” implica “ p ou r ”.

Esses são os princípios *formais* da dedução empregados em *Principia Mathematica*. Um princípio formal de dedução tem um duplo uso, e é para deixar isso claro que citamos as cinco proposições anteriores. Tem um uso como a premissa de uma inferência e um uso ao estabelecer o fato de que a premissa implica uma conclusão. No esquema de uma inferência temos uma proposição p , e uma proposição “ p implica q ”, da qual inferimos q . Ora, quando estamos interessados nos princípios de dedução, nosso aparato de proposições primitivas deve produzir tanto o p quanto o “ p implica q ” de nossas inferências. Isto é, nossas regras de dedução devem ser usadas *não somente como regras*, que é seu uso para estabelecer que “ p implica q ”, mas *também* como premissas substantivas, isto é, como o p de nosso esquema. Suponhamos, por exemplo, que desejemos provar que se p implica q , então se q implica r , segue-se que p implica r . Temos aqui uma relação de três proposições que afirmam implicações. Ponhamos

$$p_1 = p \text{ implica } q, p_2 = q \text{ implica } r, \text{ e } p_3 = p \text{ implica } r.$$

Temos então de provar que p_1 implica que p_2 implica p_3 . Agora tomemos o quinto de nossos princípios citados, substituamos p por não- p , lembrando que “não- p ou q ” é o mesmo que “ p implica q ”. Assim, nosso quinto princípio produz:

“Se q implica r , então ‘ p implica q ’ implica ‘ p implica r ’”, isto é, “ p_2 implica que p_1 implica p_3 .” Chamemos isso de proposição A.

Mas o quarto de nossos princípios, quando substituimos não- p e não- q por p e q , e lembramos a definição de implicação, torna-se:

“Se p implica que q implica r , então q implica que p implica r .”

Escrevendo p_2 no lugar de p , p_1 no lugar de q , e p_3 no lugar de r , isso se torna:

“Se p_2 implica que p_1 implica p_3 , então p_1 implica que p_2 implica p_3 .”
Chamemos isto B.

Provamos agora por meio de nosso quinto princípio que

“ p_2 implica que p_1 implica p_3 ”, que é o que chamamos A.

Temos assim aqui um caso do esquema de inferência, visto que A representa o p de nosso esquema, e B representa o “ p implica q ”. Chegamos, portanto, a q , a saber,

“ p implica que p_2 implica p_3 ”.

o que era a relação a ser provada. Nessa prova, a adaptação de nosso quinto princípio, que produz A, ocorre como uma premissa substantiva; ao passo que a adaptação de nosso quarto princípio, que produz B, é usada para dar a *forma* da inferência. Os empregos formais e materiais de premissas na teoria das deduções são estreitamente interligados, e não é muito importante mantê-los separados, contanto que compreendamos que são distintos em teoria.

O método mais primitivo de chegar a novos resultados a partir de uma premissa está ilustrado na dedução apresentada, mas essa mal pode ela própria ser chamada dedução. As proposições primitivas, sejam elas quais forem, devem ser consideradas afirmadas para todos os valores das proposições variáveis p , q , r que ocorrem nelas. Podemos, portanto, substituir (digamos) p por qualquer expressão cujo valor seja sempre uma proposição, por exemplo, não- p , “ s implica t ”, e assim por diante. Por meio dessas substituições obtemos realmente conjuntos de casos especiais de nossa proposição original, mas de um ponto de vista prático o que obtemos de fato são novas proposições. A legitimidade de substituições desse tipo tem de ser assegurada por meio de um princípio não-formal de inferência.⁵

Podemos agora formular o único princípio formal de inferência a que M. Nicod reduziu os cinco dados anteriores. Para esse propósito, vamos

primeiro mostrar como certas funções de verdade podem ser definidas em termos de incompatibilidade. Já vimos que

$$p \mid (q/q) \text{ significa “} p \text{ implica } q\text{”}$$

Observamos agora que

$$p \mid (q/r) \text{ significa “} p \text{ implica tanto } q \text{ quanto } r\text{”}.$$

Pois essa expressão significa “ p é incompatível com a incompatibilidade de q e r ”, isto é, “ p implica que q e r não são incompatíveis”, ou seja, “ p implica que q e r são ambos verdadeiros” — pois, como vimos, a conjunção de q e r é a negação de sua incompatibilidade.

Observemos em seguida que $t \mid (t/t)$ significa “ t implica a si mesmo”. Esse é um caso particular de $p \mid (q/q)$.

Vamos escrever p — para a negação de p assim p/s significará a negação de p/s —, isto é, significará a conjunção de p e s . Segue-se que

$$(s/q) \mid p/s —$$

expressa a incompatibilidade de s/q com a conjunção de p e s ; em outras palavras, afirma que se p e s forem ambos verdadeiros, s/q é falso, ou seja, s e q são ambos verdadeiros; em palavras ainda mais simples, afirma que p e s conjuntamente implicam s e q conjuntamente.

Agora, ponhamos $P = p \mid (q/r)$,

$$\pi = t \mid (t/t)$$

$$Q = (s/q) \mid p/s —.$$

O único princípio formal de dedução de M. Nicod é então

$$P \mid \pi / Q$$

em outras palavras, P implica tanto π quanto Q .

Além disso, ele emprega um princípio não-formal pertencente à teoria dos tipos (que não precisa nos preocupar), e um correspondente ao princípio segundo o qual, dado p , e dado que p implica q , podemos afirmar q . Esse princípio é

“Se $p \mid (r/q)$ for verdadeiro, e p for verdadeiro, então q é verdadeiro.”

Desse aparato segue-se toda a teoria da dedução, exceto na medida em que estamos interessados na dedução a partir ou para a existência da verdade universal de “funções proposicionais”, que consideraremos no próximo capítulo.

Há, se não me engano, certa confusão nas mentes de alguns autores quanto à relação entre proposições em virtude da qual uma inferência é válida. Para que possa ser *válido* inferir q de p , é necessário apenas que p seja verdadeiro e que a proposição “não- p ou q ” seja verdadeira. Sempre que esse for o caso, está claro que q deve ser verdadeiro. Mas a inferência só ocorrerá de fato quando a proposição “não- p ou q ” for conhecida de outra maneira que não por meio do conhecimento de não- p ou do conhecimento de q . Sempre que p for falso, “não- p ou q ” será verdadeiro, mas é inútil para a inferência, que requer que p seja verdadeiro. Sempre que já sabemos que q é verdadeiro, sabemos também, é claro, que “não- p ou q ” é verdadeiro, mas novamente isso é inútil para a inferência, uma vez que q já é conhecido e, portanto, não precisa ser inferido. De fato, a inferência só surge quando “não- p ou q ” podem ser conhecidos sem que já saibamos qual das duas alternativas torna a disjunção verdadeira. Ora, as circunstâncias em que isso ocorre são aquelas em que existem certas relações de forma entre p e q . Por exemplo, sabemos que se r implica a negação de s , então s implica a negação de r . Entre “ r implica não- s ” e “ s implica não- r ”, há uma relação formal que nos permite *saber* que a primeira implica a segunda, sem precisar saber primeiro que a primeira é falsa ou que a segunda é verdadeira. É nessas circunstâncias que a relação de implicação tem utilidade prática para se extrair inferências.

Mas essa relação formal só é requerida para que possamos ser capazes de *saber* que ou a premissa é falsa ou a conclusão é verdadeira. É a verdade de “não- p ou q ” que é requerida para a *validade* da inferência; o que é requerido adicionalmente é requerido apenas para a exeqüibilidade prática da inferência. O professor C.I. Lewis⁶ estudou especialmente a relação mais estreita, formal, que podemos chamar “dedutibilidade formal”. Ele insiste em que a relação mais ampla, aquela expressa por “não- p ou q ”, não deveria ser chamada “implicação”. Isso é, contudo, uma questão de palavras. Contanto que nosso uso das palavras seja coerente, pouco importa como as definamos. O ponto essencial de diferença entre a teoria que eu defendo e a defendida pelo professor Lewis é este: ele

sustenta que, quando uma proposição q é “formalmente dedutível” de outra p , a relação que percebemos entre elas é uma relação que ele chama “implicação estrita”, que não é a relação expressa por “não- p ou q ”, mas uma relação mais estreita, que vigora apenas quando há certas conexões formais entre p e q . Eu sustento que, quer essa relação de que ele fala exista ou não, ela é de todo modo uma relação de que a matemática não precisa e, portanto, uma relação que, por razões gerais de economia, não deveria ser admitida em nosso aparato de noções fundamentais; que, seja qual for a relação de “dedutibilidade formal” que vigore entre duas proposições, ocorre que podemos ver ou que a primeira é falsa ou a segunda verdadeira, e que não precisamos admitir em nossas premissas nada além desse fato; e que, finalmente, as razões de detalhe que o professor Lewis aduz contra a idéia que defendo podem ser todas refutadas em detalhe, e dependem, para sua plausibilidade, de uma suposição oculta e inconsciente do ponto de vista que rejeito. Concluo, portanto, que não há nenhuma necessidade de admitir como noção fundamental nenhuma forma ou implicação não exprimível como função de verdade.

Capítulo 15

Funções proposicionais

Quando discutimos proposições, no capítulo anterior, não tentamos dar uma definição da palavra “proposição”. Embora ela não possa ser formalmente definida, é necessário dizer alguma coisa sobre seu significado, de modo que evite a confusão muito comum com “funções proposicionais”, que serão o tópico desse capítulo.

Entendemos principalmente por “proposição” uma forma de palavras que expressa o que é ou verdadeiro ou falso. Digo “principalmente” porque não desejo excluir outros símbolos que não os verbais, ou até meros pensamentos, se tiverem um caráter simbólico. Mas penso que a palavra “proposição” deveria ser limitada ao que pode, em algum sentido, ser chamado “símbolos” e, ademais, a símbolos tais que dêem expressão a verdade e falsidade. Desse modo, “dois e dois são quatro” e “dois e dois são cinco” serão proposições, e assim também será “Sócrates é um homem” e “Sócrates não é um homem”. A afirmação: “Não importa que números a e b possam ser, $(a + b)^2 = a^2 + 2ab + b^2$ ” é uma proposição; mas a simples fórmula “ $(a + b)^2 = a^2 + 2ab + b^2$ ” sozinha não é, já que nada afirma de definido, a menos que, ademais, sejamos informados, ou levados a supor, que a e b devem ter todos os valores possíveis, ou devem ter tais e tais valores. A primeira dessas coisas é tacitamente admitida, em regra, na enunciação de fórmulas matemáticas, que assim se tornam proposições; mas se nenhuma admissão desse tipo fosse feita, elas seriam “funções proposicionais”. Uma “função proposicional”, de fato, é uma expressão que contém um ou mais constituintes indeterminados, tais que, quando valores são atribuídos a esses constituintes, a expressão se torna uma proposição. Em outras palavras, é uma função cujos valores são proposições. Mas essa última definição deve ser usada com cautela. Uma função descritiva, por exemplo, “a mais difícil proposição no tratado de

matemática de A ”, não será uma função proposicional, embora seus valores sejam proposições. Nesse caso, porém, as proposições são apenas descritas: numa função proposicional, os valores devem realmente *enunciar* proposições.

É fácil dar exemplos de funções proposicionais: “ x é humano” é uma função proposicional; enquanto x permanecer indeterminado, não é nem verdadeira nem falsa, mas quando um valor for atribuído a x , se tornará uma proposição verdadeira ou falsa. Toda equação matemática é uma função proposicional. Enquanto as variáveis não tiverem valor definido, a equação é meramente uma expressão à espera de determinação para se tornar uma proposição verdadeira ou falsa. Se for uma equação contendo uma variável, torna-se verdadeira quando a variável é igualada a uma raiz da equação; de outro modo torna-se falsa; mas se for uma “identidade”, será verdadeira quando a variável for qualquer número. A equação para uma curva num plano ou para uma superfície no espaço é uma função proposicional, verdadeira para valores das coordenadas que pertencem a pontos na curva ou na superfície, falsa para outros valores. Expressões de lógica tradicional como “todo A é B ” são funções proposicionais: A e B têm de ser determinados como classes definidas antes que tais expressões se tornem verdadeiras ou falsas.

A noção de “casos” ou “exemplos” depende de funções proposicionais. Considere, por exemplo, o tipo de processo sugerido pelo que é chamado “generalização”, e tomemos um exemplo muito primitivo, digamos “o relâmpago é seguido pelo trovão”. Temos uma série de “casos” disso, isto é, um número de proposições tais como: “Isso é um relâmpago e ele é seguido por um trovão.” Essas ocorrências são “casos” de quê? São casos da função proposicional: “Se x é um relâmpago, x é seguido por um trovão.” O processo de generalização (em cuja validade felizmente não estamos interessados) consiste em passar de um número desses casos à verdade *universal* da função proposicional: “Se x é um relâmpago, x é seguido por um trovão.” Verificaremos que, de maneira análoga, funções proposicionais estão sempre envolvidas quando quer que falemos de casos ou exemplos.

Não precisamos perguntar “Que é uma função proposicional?”, nem tentar responder a essa pergunta. Uma função proposicional isolada pode ser tomada por um mero esquema, uma mera casca, um receptáculo vazio para significado, não algo já significativo. Estamos interessados em

funções proposicionais num sentido amplo de duas maneiras: primeiro, como envolvidas nas noções “verdadeiro em todos os casos” e “verdadeiro em alguns casos”; segundo, como envolvidas na teoria das classes e relações. Adiaremos o segundo desses tópicos para um capítulo posterior; o primeiro deverá nos ocupar agora.

Quando dizemos que algo é “sempre verdadeiro” ou “verdadeiro em todos os casos” é claro que o “algo” envolvido não pode ser uma proposição. Uma proposição é simplesmente verdadeira ou falsa, e mais nada. Não há exemplos ou casos de “Sócrates é um homem” ou “Napoleão morreu em Santa Helena”. Essas são proposições, e seria sem sentido dizer que são verdadeiras “em todos os casos”. Essa expressão só é aplicável a *funções* proposicionais. Tomemos, por exemplo, o que se diz muitas vezes quando se discute causação. (Não estamos interessados na verdade ou falsidade do que é dito, mas unicamente em sua análise lógica.) Dizemos que A é, em todos os casos, seguido por B. Ora, se há “casos” de A, A deve ser algum conceito geral sobre o qual é significativo dizer “ x_1 é A”, “ x_2 é A”, “ x_3 é A”, e assim por diante, em que x_1, x_2, x_3 são particulares não idênticos uns aos outros. Isso se aplica, por exemplo, a nosso caso anterior do relâmpago. Dizemos que relâmpago (A) é seguido por trovão (B). Mas os diferentes clarões são particulares, não idênticos, embora partilhem a propriedade comum de serem relâmpagos. Geralmente a única maneira de expressar uma propriedade comum é dizer que uma propriedade comum de vários objetos é uma função proposicional que se torna verdadeira quando qualquer um desses objetos é tomado como o valor da variável. Aqui, todos os objetos são “casos” da verdade da função proposicional — pois uma função proposicional, embora não possa ser ela mesma verdadeira ou falsa, é verdadeira em certos casos e falsa em outros, a menos que seja “sempre verdadeira” ou “sempre falsa”. Quando, para retornar a nosso exemplo, dizemos que A é em todos os casos seguido por B, queremos dizer que, qualquer que seja x , se x for um A ele será seguido por B; isto é, estamos afirmando que certa função proposicional é “sempre verdadeira”.

A interpretação de sentenças que envolvam palavras como “todos”, “cada”, “um”, “o”, “alguns” requer funções proposicionais. O modo como funções proposicionais ocorrem pode ser explicado por meio de duas das palavras citadas, a saber, “todos” e “alguns”.

Em última análise, apenas duas coisas podem ser feitas com uma função proposicional: uma é afirmar que é verdadeira em *todos* os casos; a outra é afirmar que é verdadeira em pelo menos um caso, ou em *alguns* casos (como diremos, supondo que não deva haver nenhuma implicação necessária de uma pluralidade de casos). Todos os outros usos das funções proposicionais podem ser reduzidos a esses dois. Quando dizemos que uma função proposicional é verdadeira “em todos os casos”, ou “sempre” (como diremos também, sem nenhuma sugestão temporal), queremos dizer que todos os seus valores são verdadeiros. Se “ ϕx ” for a função, e a o tipo correto de objeto para ser um argumento para “ ϕx ”, então ϕa deve ser verdadeiro, não importa como a tenha sido escolhido. Por exemplo, “se a é humano, a é mortal” é verdadeiro quer a seja humano ou não; de fato, toda proposição dessa forma é verdadeira. Assim, a função proposicional “se x é humano, x é mortal” é “sempre verdadeira” ou “verdadeira em todos os casos”. Ora, a afirmação “não existem unicórnios” é o mesmo que a afirmação “a função proposicional ‘ x não é um unicórnio’ é verdadeira em todos os casos”. As asserções feitas no capítulo anterior sobre proposições, por exemplo, “ p ou q implica ‘ q ou p ’”, são na realidade asserções de que certas funções proposicionais são verdadeiras em todos os casos. Não afirmamos que esse princípio, por exemplo, é verdadeiro somente acerca desse ou daquele p ou q particular, mas que é verdadeiro acerca de *qualquer* p ou q com relação ao qual possa ser formulado de maneira significativa. A condição de que uma função deve ser *significativa* para um dado argumento é igual à condição de que deve ter um valor para aquele argumento, seja verdadeiro ou falso. O estudo das condições de significação pertence à doutrina dos tipos, que não exploraremos além do esboço dado no capítulo anterior.

Não só os princípios da dedução, mas todas as proposições primitivas da lógica, consistem em asserções de que certas funções proposicionais são sempre verdadeiras. Se esse não fosse o caso, elas teriam de mencionar coisas ou conceitos particulares — Sócrates ou vermelhidão, ou leste e oeste, ou seja o que for —, e claramente não é da competência da lógica fazer asserções verdadeiras com relação a tal coisa ou conceito, mas não com relação a tal outra. É parte da definição da lógica (mas não a totalidade de sua definição) que todas as suas proposições são completamente gerais, isto é, todas consistem na asserção de que uma função proposicional que não contém nenhum termo constante é sempre

verdadeira. Retornaremos em nosso capítulo final à discussão de funções proposicionais que não contêm nenhum termo constante. Por ora, passaremos à outra coisa que pode ser feita com uma função proposicional, a saber, a asserção de que ela “às vezes é verdadeira”, isto é, verdadeira em pelo menos um caso.

Quando dizemos “há homens”, isso significa que a função proposicional “ x é um homem” às vezes é verdadeira. Quando dizemos “alguns homens são gregos”, isso significa que a função proposicional “ x é homem e grego” às vezes é verdadeira. Quando dizemos “canibais ainda existem na África”, isso significa que a função proposicional “ x é um canibal atualmente na África” é às vezes verdadeira, ou seja, é verdadeira para alguns valores de x . Dizer “há pelo menos n indivíduos no mundo” é dizer que a função proposicional “ α é uma classe de indivíduos e membro do número cardinal n ” às vezes é verdadeira, ou, como também poderíamos dizer, é verdadeira para certos valores de α . Essa forma de expressão é mais conveniente quando é necessário indicar qual é o constituinte variável que estamos tomando como argumento para nossa função proposicional. Por exemplo, a função proposicional citada, que podemos abreviar para “ α é uma classe de n indivíduos”, contém duas variáveis, α e n . O axioma da infinidade, na linguagem das funções proposicionais, é: “A função proposicional ‘se n for um número indutivo, será verdadeiro para alguns valores de α que α é uma classe de n indivíduos’ será verdadeira para todos os valores possíveis de n .” Aqui há uma função subordinada, “ α é uma classe de n indivíduos”, que dizemos ser *às vezes* verdadeira com relação a α ; e dizemos que a asserção de que isso acontece se n for um número indutivo é sempre verdadeira com relação a n .

A afirmação de que uma função ϕx é sempre verdadeira é a negação da afirmação que não- ϕx às vezes é verdadeira, e a afirmação de que ϕx às vezes é verdadeira é a negação da afirmação de que não- ϕx é sempre verdadeira. Assim, a afirmação “todos os homens são mortais” é a negação da afirmação de que a função “ x é um homem imortal” às vezes é verdadeira. E a afirmação de que “existem unicórnios” é a negação da afirmação de que a função “ x não é um unicórnio” é sempre verdadeira.¹ Dizemos que ϕx “nunca é verdadeiro” ou “é sempre falso” se não- ϕx for sempre verdadeiro. Podemos, se quisermos, tomar um do par “‘sempre’, ‘às vezes’” como uma idéia primitiva, e definir a outra por meio de uma e da negação. Assim, se escolhermos “às vezes” como nossa idéia primitiva,

podemos definir: “‘ $\forall x$ é sempre verdadeiro’ deve significar ‘é falso que não- $\forall x$ seja às vezes verdadeiro’”.² Por razões ligadas à teoria dos tipos, porém, parece mais correto tomar “sempre” e “às vezes” como idéias primitivas e definir por meio delas a negação de proposições em que ocorrem. Isto é, supondo que já definimos (ou adotamos como idéia primitiva) a negação de proposições do tipo a que x pertence, definimos: “A negação de “ $\forall x$ sempre” é “não- $\forall x$ às vezes”; e a negação de “ $\forall x$ às vezes” é “não- $\forall x$ sempre”. De maneira semelhante, podemos redefinir disjunção e outras funções de verdade, tal como aplicadas a proposições que contenham variáveis aparentes, em termos das definições e idéias primitivas para proposições que não contenham variáveis aparentes. Proposições que não contêm variáveis aparentes são chamadas “proposições elementares”. A partir dessas podemos subir passo a passo, usando os métodos que acabamos de indicar, até a teoria das funções de verdade tal como aplicada a proposições contendo uma, duas, três . . . variáveis, ou qualquer número até n , em que n é um número finito designado.

As formas consideradas as mais simples na lógica formal tradicional, estão na realidade longe de o serem, e todas envolvem a afirmação de todos os valores ou de alguns valores de uma função proposicional composta. Tomemos, para começar, “todo S é P”. Consideraremos que S é definido por uma função proposicional $\forall x$, e P por uma função proposicional ψx . Por exemplo, se S for *homens*, $\forall x$ será “ x é humano”; se P for *mortais*, ψx será “há um momento em que x morre”. Então “todo S é P” significa: “‘ $\forall x$ implica ψ ’ é sempre verdadeiro”. Deve-se observar que “todo S é P” não se aplica apenas àqueles termos que realmente são S’s; diz alguma coisa igualmente sobre termos que não são S’s. Suponhamos que deparamos com um x sobre o qual não sabemos se é ou não um S; apesar disso, nossa afirmação “todo S é P” nos diz alguma coisa sobre x , a saber, que se x for um S, então x será um P. E isso é tão verdadeiro quer x seja um S ou não. Se não fosse igualmente verdadeiro em ambos os casos, a *reductio ad absurdum* não seria um método válido; pois a essência desse método consiste em usar implicações em casos em que (como mais tarde se revelará) a hipótese é falsa. Podemos expressar isso de outra maneira. Para compreender “todo S é P”, não é necessário ser capaz de enumerar que termos são S’s; contanto que saibamos o que se entende por um S e o que se entende por um P, podemos entender perfeitamente o que é

realmente afirmado por “todo S é P”, por menos casos reais de ambos que possamos conhecer. Isso mostra que não são meramente os termos reais que são S’s que são relevantes na afirmação “todo S é P”, mas todos os termos com relação aos quais a suposição de que eles são S’s é significativa, isto é, todos os termos que são S’s, juntamente com todos os termos que não são S’s — ou seja, a totalidade do “tipo” lógico apropriado. O que se aplica a afirmações sobre *todo* se aplica também a afirmações sobre *alguns*. “Existem homens”, por exemplo, significa que “ x é humano” é verdadeiro para alguns valores de x . Aqui *todos* os valores de x (isto é, todos os valores de x para os quais “ x é humano” é significativo, quer verdadeiro ou falso) são relevantes, e não somente aqueles que de fato são humanos. (Isto se torna óbvio se consideramos como poderíamos provar que uma afirmação como essa é *falsa*.) Toda asserção sobre “todos” ou “alguns” envolve, portanto, não apenas os argumentos que tornam certa função verdadeira, mas todos os que a tornam significativa, isto é, todos para os quais ela tem de algum modo um valor, quer verdadeiro ou falso.

Podemos agora prosseguir com nossa interpretação das formas tradicionais da antiquada lógica formal. Supomos que S são aqueles termos x para os quais ϕx é verdadeiro, e P são aqueles para os quais ψx é verdadeiro. (Como veremos num capítulo posterior, todas as classes são derivadas dessa maneira de funções proposicionais.) Portanto:

“Todo S é P” significa “‘ ϕx implica ψx ’ é sempre verdadeiro”.

“Algum S é P” significa “‘ ϕx implica ψx ’ às vezes é verdadeiro”.

“Nenhum S é P” significa “‘ ϕx implica não- ψx ’ é sempre verdadeiro”.

“Algum S é não P” significa “‘ ϕx e não- ψx ’ às vezes é verdadeiro”.

Convém observar que as funções proposicionais afirmadas aqui para todos ou alguns valores não são ϕx e ψx eles mesmos, mas funções de verdade de ϕx e ψx para o mesmo argumento x . A maneira mais fácil de fazer uma idéia do tipo de coisa que se quer dizer é começar não de ϕx e ψx em geral, mas de ϕa e ψa , onde a é uma constante. Suponhamos que estamos considerando todos os “homens são mortais”; começaremos com

“Se Sócrates é humano, Sócrates é mortal,”

e em seguida consideraremos que Sócrates é substituído por uma variável x sempre que “Sócrates” ocorrer. O objetivo a ser assegurado é que, embora x permaneça variável, sem nenhum valor definido, terá o mesmo valor tanto em “ $\emptyset x$ ” como em “ νx ” quando estamos afirmando que “ $\emptyset x$ ” implica “ νx ” sempre verdade. Isso requer que comecemos com uma função cujos valores sejam tais que “ $\emptyset a$ implica νa ”, e não com duas funções separadas $\emptyset x$ e νx ; pois se começarmos com duas funções separadas, nunca poderemos assegurar que o x , embora permanecendo indeterminado, deve ter o mesmo valor em ambas.

Para abreviar, dizer “ $\emptyset x$ sempre implica νx ” quando queremos dizer “ $\emptyset x$ implica νx ” é sempre verdadeiro. Proposições da forma “ $\emptyset x$ sempre implica νx ” são chamadas “implicações formais”; esse nome é dado igualmente se houver diversas variáveis.

As definições anteriores mostram o quanto proposições como “todo S é P” estão distantes das formas mais simples com que a lógica tradicional começa. É típico da falta de análise envolvida que a lógica tradicional trate “todo S é P” como uma proposição da mesma forma que “ x é P” — por exemplo, trata “todos os homens são mortais” da mesma forma que “Sócrates é mortal”. Como acabamos de ver, a primeira é da forma “ $\emptyset x$ sempre implica νx ”, enquanto a segunda é da forma “ νx ”. A separação enfática dessas duas formas, efetuada por Peano e Frege, constituiu um avanço fundamental na lógica simbólica.

Veremos que “todo S é P” e “nenhum S é P” não diferem realmente em forma, exceto pela substituição de νx por não- νx , e que o mesmo se aplica a “algum S é P” e “algum S é não P”. É preciso observar também que as regras tradicionais de conversão são defeituosas, se adotarmos a idéia, que é a única tecnicamente tolerável, de que proposições como “todo S é P” não envolvem a “existência” de S’s, isto é, não requerem que haja termos que são S’s. Essas definições levam ao resultado de que, se $\emptyset x$ for sempre falso, ou seja, se não houver nenhum S, então “todo S é P” e “nenhum S é P” serão ambos verdadeiros, qualquer que seja P. Pois, segundo a definição dada no capítulo anterior, “ $\emptyset x$ implica νx ” significa “não- $\emptyset x$ ou νx ”, o que é sempre verdadeiro se não- $\emptyset x$ for sempre verdadeiro. A princípio esse resultado poderia levar o leitor a desejar definições diferentes, mas um pouco de experiência prática logo mostra que quaisquer definições diferentes seriam inconvenientes e ocultariam as idéias importantes. A proposição “ $\emptyset x$ sempre implica ψx , e $\emptyset x$ às vezes é verdadeiro” é

essencialmente compósita, e seria muito desajeitado dar isso como a definição de “todo S é P”, pois nesse caso não nos sobraria linguagem para “ $\forall x$ sempre implica $\forall x$ ”, que é cem vezes mais necessária que a outra. Mas, com nossa definição, “todo S é P” não implica “algum S é P”, visto que a primeira permite a não-existência de S, e a segunda, não; assim, conversão *per accidens* torna-se inválida, e alguns modos do silogismo são falaciosos, por exemplo, Darapti: “Todo M é S, todo M é P, portanto algum S é P”, que falha se não houver nenhum M.

A noção de “existência” tem várias formas, uma das quais nos ocupará no próximo capítulo; mas a forma fundamental é aquela derivada imediatamente da noção de “às vezes verdadeiro”. Dizemos que um argumento *a* “satisfaz” uma função $\forall x$ se $\forall a$ for verdadeira; é nesse mesmo sentido que dizemos que as raízes de uma equação satisfazem a equação. Ora, se $\forall x$ às vezes é verdadeiro, podemos dizer que há x 's para o quais ela é verdadeira, ou podemos dizer “existem argumentos que satisfazem $\forall x$ ”. Esse é o significado fundamental da palavra “existência”. Outros significados ou são derivados desse ou envolvem alguma confusão de pensamento. Podemos dizer corretamente “existem homens”, significando que “ x é um homem”, às vezes é verdadeiro. Mas se fizermos um pseudo-silogismo: “Existem homens, Sócrates é um homem, logo existe Sócrates”, estaremos dizendo um absurdo, pois “Sócrates” não é, como “homens” meramente um argumento indeterminado para uma função proposicional dada. A falácia é estreitamente análoga àquela do raciocínio: “Os homens são numerosos, Sócrates é um homem, logo Sócrates é numeroso.” Nesse caso é óbvio que a conclusão é absurda, mas no caso da existência não é óbvio, por razões que aparecerão de maneira mais completa no próximo capítulo. Por enquanto, contentemo-nos em observar o fato de que, embora seja correto dizer “existem homens”, é incorreto, ou melhor, sem sentido, atribuir existência a um x particular dado que por acaso é um homem. Em geral, “existem termos que satisfazem $\forall x$ ” significa “ $\forall x$ às vezes é verdade”; mas “existe *a*” (onde *a* é um termo que satisfaz $\forall x$) é um mero ruído ou formato, desprovido de significação. Veremos que, tendo em mente essa simples falácia, podemos resolver muitos enigmas filosóficos antigos concernentes ao significado da existência.

Um outro conjunto de noções com relação ao qual a filosofia se permitiu cair em insanáveis confusões por não distinguir suficientemente

proposições e funções proposicionais são as noções de “modalidade”: *necessário*, *possível* e *impossível*. (Às vezes *contingente* ou *assertórico* é usado em vez de possível.) A idéia tradicional era que, entre proposições verdadeiras, algumas eram necessárias, enquanto outras eram verdadeiras por mero acaso. Contudo, nunca houve uma explicação clara do que era acrescentado à verdade pela concepção de necessidade. No caso das funções proposicionais, a tríplice divisão é óbvia. Se “ ϕx ” for um valor indeterminado de certa função proposicional, ele será *necessário* se a função sempre for verdadeira, *possível* se ela for verdadeira às vezes e *impossível* se ela nunca for verdadeira. Esse tipo de situação surge em relação à probabilidade, por exemplo. Suponhamos que uma bola x seja retirada de um saco que contém certo número de bolas: se todas as bolas forem brancas, “ x é branca” é necessário; se algumas forem brancas, “ x é branca” é possível; se nenhuma, é impossível. Aqui, só o que se *sabe* sobre x é que ela satisfaz uma certa função proposicional, a saber “ x é uma bola que estava no saco”. Esta é uma situação geral em problemas de probabilidade e não incomum na vida prática — por exemplo, quando somos procurados por uma pessoa sobre a qual não sabemos nada, a não ser que ela traz uma carta de apresentação de nosso amigo fulano de tal. Em todos esses casos, como em relação à modalidade em geral, a função proposicional é relevante. Para um pensamento claro, em muitas diferentes direções, o hábito de manter funções proposicionais nitidamente separadas de proposições é da máxima importância, e a incapacidade de fazê-lo no passado foi uma desgraça para a filosofia.

Capítulo 16

Descrições

Tratamos no capítulo anterior das palavras *todos* e *alguns*; neste capítulo vamos considerar o artigo *o* [ou *a*] e no próximo capítulo vamos considerar o mesmo artigo no plural, *os* [ou *as*]. Pode-se julgar excessivo dedicar dois capítulos a uma palavra, mas para o matemático filósofo é uma palavra de grande importância, como o gramático de Browning com o enclítico $\delta\epsilon$, eu daria a doutrina dessa palavra se estivesse “morto da cintura para baixo” e não meramente numa prisão.

Já tivemos ocasião de mencionar “funções descritivas”, isto é, expressões como “o pai de x ” ou “o seno de x ”. Para defini-las, é preciso antes definir “descrições”.

Uma “descrição” pode ser de dois tipos, definida e indefinida (ou ambígua). Uma descrição indefinida é uma expressão da forma “uma tal coisa”, e uma descrição definida é uma expressão da forma “a tal coisa” (no singular). Começemos com a primeira.

“Quem você encontrou?” “Encontrei um homem.” “Essa é uma descrição muito indefinida.” Não estamos, portanto, nos afastando do uso em nossa terminologia. Nossa questão é: que digo eu realmente quando afirmo “Encontrei um homem”? Suponhamos, por enquanto, que minha asserção seja verdadeira e que de fato encontrei Jones. É claro que o que afirmo *não* é “Encontrei Jones”. Posso dizer “Encontrei um homem, mas não era Jones”; nesse caso, embora eu minta, não me contradigo, como o faria se, quando digo que encontrei um homem, quisesse realmente dizer que encontrei Jones. É claro que a pessoa com quem estou falando pode entender o que eu digo, mesmo que seja um estrangeiro que nunca ouviu falar de Jones.

No entanto, podemos ir mais longe: não só Jones, mas nenhum homem real, entra em minha afirmação. Isso se torna óbvio quando a afirmação é falsa, visto que não há mais razão para se supor que Jones devesse entrar

na proposição que qualquer outra pessoa. Na verdade, a afirmação permaneceria significativa, embora não tivesse possibilidade de ser verdade, mesmo que não houvesse homem nenhum. “Encontrei um unicórnio” ou “encontrei uma serpente marinha” é uma asserção perfeitamente significativa, caso saibamos o que vem a ser um unicórnio ou uma serpente marinha, isto é, qual é a definição desses monstros fabulosos. Assim, é somente o que podemos chamar de *conceito* que entra na proposição. No caso de “unicórnio”, por exemplo, há apenas o conceito: não há além disso, em algum lugar entre as sombras, algo irreal que poderia ser chamado “um unicórnio”. Portanto, como é significativo (embora falso) dizer “Encontrei um unicórnio”, é claro que essa proposição, corretamente analisada, não contém um constituinte “um unicórnio”, embora contenha o conceito “unicórnio”.

A questão da “irrealidade”, com que nos defrontamos neste ponto, é muito importante. Induzidos em erro pela gramática, a grande maioria dos lógicos que lidaram com essa questão o fez em linhas enganosas. Considerou a forma gramatical como um guia mais seguro na análise do que de fato é. E não souberam distinguir que diferenças na forma gramatical são importantes. Tradicionalmente, “Encontrei Jones” e “Encontrei um homem” contariam como proposições da mesma forma, mas na realidade são de formas muito diferentes: a primeira nomeia uma pessoa real, Jones; ao passo que a segunda envolve uma função proposicional, e se torna, quando explicitada: “A função ‘Encontrei x e x é humano’ é verdade às vezes.” (Lembremos que adotamos a convenção de usar “às vezes” como não implicando mais de uma vez.) Essa proposição obviamente não é da forma “encontrei x ”, que justifica a existência da proposição “Encontrei um unicórnio” apesar do fato de não existir tal coisa como “um unicórnio”.

Por falta do aparato das funções proposicionais, muitos lógicos foram levados à conclusão de que há objetos irrealis. Há quem afirme, por exemplo, Meinong,¹ que podemos falar sobre “a montanha dourada”, “o quadrado redondo”, e assim por diante; podemos fazer proposições verdadeiras das quais esses são os sujeitos; portanto, eles devem ter algum tipo de existência lógica, visto que de outro modo as proposições em que ocorrem seriam sem sentido. Em tais teorias, parece-me, há uma falha daquele senso de realidade que deveria ser preservado mesmo nos estudos mais abstratos. A lógica, eu afirmaria, não deve admitir unicórnios mais

do que a zoologia pode admiti-los; pois a lógica diz respeito ao mundo real tão verdadeiramente quanto a zoologia, embora com suas características mais abstratas e gerais. Dizer que unicórnios têm existência na heráldica, ou na literatura, ou na imaginação, é uma evasiva das mais desprezíveis e mesquinhas. O que existe na heráldica não é um animal, de carne e osso, que se move e respira por iniciativa própria. O que existe é uma imagem, ou uma descrição em palavras. De maneira semelhante, sustentar que Hamlet, por exemplo, existe em seu próprio mundo, a saber, o mundo da imaginação de Shakespeare, tão verdadeiramente quanto (digamos) Napoleão existiu no mundo comum, é dizer algo que confunde deliberadamente, ou é confuso num grau quase inacreditável. Há somente um mundo, o mundo “real”: a imaginação de Shakespeare é parte desse mundo, e os pensamentos que ele teve ao escrever Hamlet são reais. Assim também são os pensamentos que temos ao ler a peça. Mas é da própria essência da ficção que somente os pensamentos, sentimentos etc. em Shakespeare e em seus leitores sejam reais, e que não haja, em adição a eles, nenhum Hamlet objetivo. Quando levamos em conta todos os sentimentos despertados por Napoleão em escritores e leitores de História, não tocamos no homem real; mas no caso de Hamlet, ao fazer isso nós o esgotamos. Se ninguém pensasse sobre Hamlet, não sobraria nada dele; se ninguém tivesse pensado sobre Napoleão, ele logo trataria de forçar alguém a fazê-lo. O senso de realidade é vital em lógica, e quem quer que faça malabarismos com ele, alegando que Hamlet tem outro tipo de realidade, está fazendo um desserviço ao pensamento. Um senso de realidade robusto é muito necessário para se formular uma análise correta de proposições sobre unicórnios, montanhas douradas, quadrados redondos e outros desses pseudo-objetos.

Em obediência ao senso de realidade, devemos insistir em que, na análise de proposições, nada “irreal” deve ser admitido. Mas, afinal de contas, se não *há* nada irreal, como *poderíamos* — pode-se perguntar — admitir alguma coisa irreal? A resposta é que, ao lidar com proposições, estamos lidando em primeira instância com símbolos, e se atribuímos significado a grupos de símbolos que não têm significado, cairemos no erro de admitir irrealidades, no único sentido em que isso é possível, a saber, como objetos descritos. Na proposição “Encontrei um unicórnio”, as três palavras juntas compõem uma proposição significativa, exatamente no mesmo sentido que a palavra “homem”. Mas as *duas* palavras “um

unicórnio” não formam um grupo subordinado que tenha um sentido próprio. Assim, se atribuímos falsamente sentido a essas duas palavras, vemo-nos às voltas com “um unicórnio”, e com o problema de como pode haver tal coisa num mundo em que não há unicórnios. “Um unicórnio” é uma descrição indefinida que não descreve nada. Não é uma descrição indefinida que descreve algo irreal. Uma proposição como “ x é irreal” só tem sentido quando “ x ” é uma descrição, definida ou indefinida; nesse caso a proposição será verdadeira se “ x ” for uma descrição que não descreve nada. Mas quer a descrição “ x ” descreva alguma coisa ou não descreva nada, de todo modo ela não é um constituinte da proposição em que ocorre; como “um unicórnio” há pouco, ela não é um grupo subordinado dotado de um sentido próprio. Tudo isso resulta do fato de que, quando “ x ” é uma descrição, “ x é irreal” ou “ x não existe” não é absurdo, mas é sempre significativo e às vezes verdadeiro.

Podemos agora seguir adiante para definir em geral o sentido de proposições que contêm descrições ambíguas. Suponhamos que desejamos fazer alguma afirmação sobre “uma tal coisa”, em que “tais coisas” são aqueles objetos que têm certa propriedade $\emptyset$, isto é, aqueles objetos x para os quais a função $\emptyset x$ é verdadeira. (Por exemplo, se tomarmos “um homem” como nosso exemplo de “uma tal coisa”, $\emptyset x$ será “ x é humano”.) Suponhamos agora que desejamos afirmar a propriedade $\emptyset$ de “tal coisa”, isto é, queremos afirmar que “uma tal coisa” possui aquela propriedade que x tem quando $\emptyset x$ é verdadeira. (Por exemplo, no caso de “encontrei um homem”, $\emptyset x$ será “Encontrei x ”.) Ora, a proposição de que “uma tal coisa” tem a propriedade $\emptyset$ não é uma proposição da forma “ $\emptyset x$ ”. Se fosse, “uma tal coisa” teria de ser idêntico a x para um x adequado; e embora (num certo sentido) isso possa ser verdadeiro em alguns casos, certamente não é verdadeiro num caso tal como “um unicórnio”. É exatamente este fato, de que a afirmação de que uma tal coisa tem a propriedade $\emptyset$ não é da forma $\emptyset x$, que torna possível para “uma tal coisa” ser, num certo sentido claramente definível, “irreal”. A definição é a seguinte:

A afirmação de que “um objeto dotado da propriedade $\emptyset$ tem a propriedade $\emptyset$ ”

significa:

“A afirmação conjunta de $\emptyset x$ e $\emptyset x$ não é sempre falsa.”

Do ponto de vista da lógica, essa é a mesma proposição que poderia ser expressada por “alguns ϕ 's são ψ 's”; retoricamente, porém, há uma diferença, porque em um caso há uma sugestão de singularidade e no outro de pluralidade. Esse, contudo, não é o ponto importante. O ponto importante é que, quando corretamente analisadas, proposições que tratam verbalmente de “uma tal coisa” revelam não conter nenhum constituinte representado por essa expressão. E é por isso que tais proposições podem ser significativas mesmo quando algo como uma tal coisa não existe.

A definição de *existência*, tal como aplicada a descrições ambíguas, resulta do que foi dito no capítulo anterior. Dizemos que “existem homens” ou que “existe um homem” se a função proposicional “ x é humano” for verdadeira às vezes; e em geral “uma tal coisa” existe se “ x é uma tal coisa” for verdadeira às vezes. Podemos pôr isso numa outra linguagem. A proposição “Sócrates é um homem” é sem dúvida *equivalente* a “Sócrates é humano”, mas não se trata exatamente da mesma proposição. O *é* de “Sócrates é humano” expressa a relação entre sujeito e predicado; o *é* de “Sócrates é um homem” expressa identidade. É uma desgraça para a raça humana que ela tenha escolhido empregar a mesma palavra “é” para essas duas idéias inteiramente diferentes — uma desgraça que uma linguagem lógica simbólica obviamente consegue remediar. A identidade em “Sócrates é um homem” é identidade entre um objeto nomeado (aceitando “Sócrates” como um nome, sujeito a qualificações explicadas mais tarde) e um objeto ambigualmente descrito. Um objeto ambigualmente descrito “existirá” quando pelo menos uma proposição dessas for verdadeira, isto é, quando houver pelo menos uma proposição verdadeira da forma “ x é uma tal coisa”, em que “ x ” é um nome. É característico de descrições ambíguas (em contraposição a descrições definidas) que pode haver qualquer número de proposições verdadeiras da forma apresentada — Sócrates é um homem, Platão é um homem etc. Assim, “existe um homem” segue-se de Sócrates, ou de Platão ou de qualquer outra pessoa. Com descrições definidas, por outro lado, a forma correspondente de proposição, a saber “ x é a tal coisa” (em que x é um nome), só pode ser verdade para no máximo um valor de x . Isso nos leva ao tema das descrições definidas, que devem ser definidas de maneira análoga àquela empregada para descrições ambíguas, porém muito mais complicada.

Chegamos agora ao principal tema do presente capítulo, a saber, a definição da palavra *o*. Um ponto muito importante sobre a definição de “uma tal coisa” aplica-se igualmente a “a tal coisa”; a definição a ser procurada é uma definição de proposições em que essa expressão ocorre, não uma definição da expressão em si mesma, isoladamente. No caso de “uma tal coisa”, isso é bastante óbvio: ninguém poderia supor que “um homem” fosse um objeto definido, que poderia ser definido por si mesmo. Sócrates é um homem, Platão é um homem, Aristóteles é um homem, mas não podemos inferir que “um homem” significa o mesmo que “Sócrates” significa, e também o mesmo que “Platão” significa e também o mesmo que “Aristóteles” significa, pois esses três nomes têm diferentes significados. No entanto, quando tivermos enumerado todos os homens no mundo, não terá sobrado nada de que possamos dizer: “Esse é um homem, e não só isso, mas é *o* ‘um homem’, a entidade quintessencial que é apenas um homem indefinido sem ser ninguém em particular.” Fica, portanto, bastante claro que tudo que há no mundo é definido: se há um homem é um homem definido e não qualquer outro. Logo, não é possível encontrar no mundo uma entidade como “um homem” em contraposição a homens específicos. E assim é natural que não definamos a expressão “um homem” em si mesma, mas somente as proposições em que ela ocorre.

No caso de “a tal coisa”, isso é igualmente verdade, embora à primeira vista menos óbvio. Podemos demonstrar que esse deve ser o caso mediante uma consideração da diferença entre um *nome* e uma *descrição definida*. Tomemos a proposição “Scott é o autor de *Waverley*”. Temos aqui um nome, “Scott”, e uma descrição, “o autor de *Waverley*”, que se afirma pertencerem à mesma pessoa. A distinção entre um nome e todos os outros símbolos pode ser explicada da seguinte maneira.

Um nome é um símbolo simples cujo significado só pode ocorrer como sujeito, isto é, algo do tipo que, no Capítulo 13, definimos como um “indivíduo” ou um “particular”. E um “símbolo” “simples” é um símbolo que não tem partes que sejam símbolos. Assim, “Scott” é um símbolo simples, porque, embora tenha partes (a saber, as diferentes letras), essas partes não são símbolos. Por outro lado, “o autor de *Waverley*” não é um símbolo simples, porque as diferentes palavras que compõem a expressão são partes que são símbolos. Se, como pode ocorrer, qualquer coisa que *pareça* ser um “indivíduo” for realmente passível de análise adicional, teremos de nos contentar com o que pode ser chamado “indivíduos

relativos”, que serão termos que, em todo o contexto em questão, nunca são analisados e nunca ocorrem de outra maneira senão como sujeitos. E nesse caso teremos, correspondentemente, de nos contentar com “nomes relativos”. Do ponto de vista de nosso problema presente, a saber, a definição de descrições, esse problema, se esses são nomes absolutos ou somente nomes relativos, pode ser ignorado, visto que ele diz respeito a diferentes estágios na hierarquia dos “tipos”, ao passo que temos de comparar pares como “Scott” e “o autor de *Waverley*”, que se aplicam ambos ao mesmo objeto, e não suscitam o problema dos tipos. Podemos, portanto, por enquanto, tratar os nomes como capazes de serem absolutos; nada que teremos de dizer dependerá dessa suposição, mas o fraseio poderá ser um pouco reduzido por isso.

Temos, portanto, duas coisas para comparar: (1) um *nome*, que é um símbolo simples, que designa diretamente um indivíduo que é seu significado, e que possui esse significado por direito próprio, independentemente dos significados de todas as outras palavras; (2) uma *descrição*, que consiste em várias palavras cujos significados já estão estabelecidos, e das quais resulta o que quer que deva ser tomado como o “significado” da descrição.

Uma proposição que contém uma descrição não é idêntica ao que essa proposição se torna quando um nome é substituído, mesmo que o nome nomeie o mesmo objeto que a descrição descreve. “Scott é o autor de *Waverley*” é obviamente uma proposição diferente de “Scott é Scott”: a primeira é um fato da história literária, a segunda é um truísmo trivial. E se puséssemos qualquer outra pessoa que não Scott em lugar de “o autor de *Waverley*”, nossa proposição se tornaria falsa, e portanto certamente não seria mais a mesma proposição. Mas, pode-se dizer, nossa proposição é essencialmente da mesma forma que (digamos) “Scott é *sir* Walter”, em que se diz que dois nomes se aplicam à mesma pessoa. A resposta é, se “Scott é *sir* Walter” realmente significa “a pessoa chamada ‘Scott’ é a pessoa chamada ‘*sir* Walter’”, então os nomes estão sendo usados como descrições, isto é, o indivíduo, em vez de estar sendo nomeado, está sendo descrito como a pessoa que tem aquele nome. Essa é uma maneira pela qual nomes são freqüentemente usados na prática, e não haverá, via de regra, nada na fraseologia para mostrar se eles estão sendo usados dessa maneira ou *como* nomes. Quando um nome é usado diretamente, apenas para indicar do que estamos falando, ele não faz parte do *fato* afirmado, ou

da falsidade, caso nossa asserção seja falsa: é meramente parte do simbolismo pelo qual expressamos nosso pensamento. O que queremos expressar é algo que poderia (por exemplo) ser transposto para uma língua estrangeira; é algo para o qual as palavras reais são um veículo, mas do qual elas não fazem parte. Por outro lado, quando fazemos uma proposição sobre “a pessoa chamada ‘Scott’”, o nome real “Scott” entra no que estamos afirmando, e não meramente na linguagem usada para se fazer a asserção. Nossa proposição será diferente se a substituirmos por “a pessoa chamada ‘*sir* Walter’”. Mas enquanto estivermos usando nomes *como* nomes, quer digamos “Scott” ou digamos “*sir* Walter” é tão irrelevante para o que estamos afirmando quanto se falamos em inglês ou francês. Assim, enquanto nomes estiverem sendo usados *como* nomes, “Scott é *sir* Walter” é a mesma proposição trivial que “Scott é Scott”. Isto completa a prova de que “Scott é o autor de *Waverley*” não é a mesma proposição que resulta quando substituímos “o autor de *Waverley*” por um nome, não importa que nome seja esse.

Quando usamos uma variável, e falamos de uma função proposicional, digamos ϕx , o processo de aplicar afirmações gerais sobre x a casos particulares consistirá em substituir a letra “ x ” por um nome, supondo que ϕ é uma função que tem indivíduos por seus argumentos. Suponhamos, por exemplo, que ϕx é “sempre verdadeiro”; digamos que é a “lei de identidade”, $x = x$. Podemos, então, substituir “ x ” por qualquer nome que escolhermos, e obteremos uma proposição verdadeira. Supondo por enquanto que “Sócrates”, “Platão” e “Aristóteles” são nomes (uma suposição muito ousada), podemos inferir da lei de identidade que Sócrates é Sócrates, Platão é Platão e Aristóteles é Aristóteles. Mas cometeremos uma falácia se tentarmos inferir, sem premissas adicionais, que o autor de *Waverley* é o autor de *Waverley*. Isso resulta do que acabamos de provar: que se substituirmos “o autor de *Waverley*” por um nome numa proposição, obteremos uma proposição diferente. Isto é, aplicando o resultado a nosso presente caso: se “ x ” for um nome, “ $x = x$ ” não será a mesma proposição que “o autor de *Waverley* é o autor de *Waverley*”, não importa que nome “ x ” possa ser. Assim, do fato de que todas as proposições da forma “ $x = x$ ” são verdadeiras, não podemos inferir, sem cerimônia, que o autor de *Waverley* é o autor de *Waverley*. De fato, proposições da forma “a tal coisa é a tal coisa” não são sempre verdadeiras: é necessário que a tal coisa *exista* (um termo que será

explicado brevemente). É falso que o atual rei da França seja o atual rei da França, ou que o quadrado redondo seja o quadrado redondo. Quando substituímos um nome por uma descrição, funções proposicionais que são “sempre verdadeiras” podem se tornar falsas, se a descrição não descrever nada. Não há mistério nenhum nisso, contanto que compreendamos (o que foi provado no capítulo anterior) que quando substituímos um nome por uma descrição o resultado não é um valor da função proposicional em questão.

Estamos agora em condições de definir proposições em que uma descrição definida ocorre. A única coisa que distingue “a tal coisa” de “uma tal coisa” é a implicação de unicidade. Não podemos falar de “o habitante de Londres” porque habitar Londres é um atributo não-único. Não podemos falar sobre “o atual rei da França” porque não há nenhum; mas podemos falar sobre “o atual rei da Inglaterra”. Assim, proposições sobre “a tal coisa” sempre implicam as proposições correspondentes sobre “uma tal coisa”, com o adendo de que não há mais que uma única tal coisa. Uma proposição como “Scott é o autor de *Waverley*” não poderia ser verdadeira se *Waverley* nunca tivesse sido escrito, ou se tivesse sido escrito por várias pessoas; e tampouco poderia ser verdadeira qualquer outra proposição resultante de uma função proposicional x pela substituição de x por “o autor de *Waverley*”. Podemos dizer que “o autor de *Waverley*” significa “o valor de x para o qual ‘ x escreveu *Waverley*’ é verdadeiro”. Dessa maneira, a proposição “o autor de *Waverley* foi Scott”, por exemplo, envolve:

- (1) “ x escreveu *Waverley*” não é sempre falso;
- (2) “se x e y escreveram *Waverley*, x e y são idênticos” é sempre verdadeiro;
- (3) “se x escreveu *Waverley*, x foi Scott” é sempre verdadeiro.

Essas três proposições, traduzidas em linguagem comum, afirmam:

- (1) pelo menos uma pessoa escreveu *Waverley*;
- (2) no máximo uma pessoa escreveu *Waverley*;
- (3) quem quer que tenha escrito *Waverley* foi Scott.

Todas essas três estão implicadas por “o autor de *Waverley* foi Scott”. Inversamente, as três juntas (mas não duas delas) implicam que o autor de

Waverley foi Scott. Por isso as três juntas podem ser tomadas como definindo o que é significado pela proposição “o autor de *Waverley* foi Scott”.

Podemos simplificar um pouco essas três proposições. A primeira e a segunda juntas são equivalentes a: “Há um termo c tal que ‘ x escreveu *Waverley*’ é verdadeiro quando x é c e é falso quando x é não c .” Em outras palavras: “Há um termo c tal que ‘ x escreveu *Waverley*’ é sempre equivalente a ‘ x é c ’.” (Duas proposições são “equivalentes” quando ambas são verdadeiras ou ambas são falsas.) Temos aqui, para começar, duas funções de x , “ x escreveu *Waverley*” e “ x é c ”, e formamos uma função de c considerando a equivalência dessas duas funções de x para todos os valores de x ; em seguida afirmamos que a função resultante de c é “verdadeira às vezes”, isto é, que é verdadeira para ao menos um valor de c . (Obviamente não pode ser verdadeira para mais que um valor de c .) Essas duas condições juntas são definidas como dando o significado de “o autor de *Waverley* existe”.

Podemos agora definir “o termo que satisfaz a função $\emptyset x$ existe”. Essa é a forma geral de que vimos anteriormente um caso particular. “O autor de *Waverley*” é “o termo que satisfaz a função ‘ x escreveu *Waverley*’”. E “a tal coisa” envolverá sempre referência a alguma função proposicional, a saber, aquela que define a propriedade que faz de uma coisa uma tal coisa. Nossa definição é como se segue: “O termo que satisfaz a função $\emptyset x$ existe” significa: “Há um termo c tal que $\emptyset x$ é sempre equivalente a ‘ x é c ’.”

Para definir “o autor de *Waverley* foi Scott”, temos ainda de levar em conta a terceira de nossas três proposições, a saber, “quem quer que tenha escrito *Waverley* foi Scott”. Isso será satisfeito meramente acrescentando-se que o c em questão deve ser Scott. Assim, “o autor de *Waverley* foi Scott” é:

“Há um termo c tal que (1) ‘ x escreveu *Waverley*’ é sempre equivalente a ‘ x é c ’, (2) c é Scott.

E de maneira geral: “o termo que satisfaz $\emptyset x$ satisfaz ux ” é definido como significando:

“Há um termo c tal que (1) $\emptyset x$ é sempre equivalente a ‘ x é c ’, (2) vc verdadeiro.”

Essa é a definição de proposições em que descrições ocorrem.

É possível ter muito conhecimento em relação a um termo descrito, isto é, conhecer muitas proposições concernentes a “a tal coisa”, sem realmente saber o que a tal coisa é, isto é, sem conhecer nenhuma proposição da forma “ x é a tal coisa” em que “ x ” é um nome. Numa história policial, proposições sobre “o homem que cometeu o ato” se acumulam, na esperança de que por fim serão suficientes para demonstrar que foi A que cometeu o ato. Podemos até chegar ao ponto de dizer que, em todo aquele conhecimento que pode ser expresso em palavras — com exceção de “este” e “aquele” e algumas outras palavras cujo significado varia em diferentes ocasiões —, não ocorre nenhum nome no sentido estrito, mas o que parece ser nome é realmente descrição. Podemos indagar, de maneira significativa, se Homero existiu, o que não poderíamos fazer se “Homero” fosse um nome. A proposição “a tal coisa existe” é significante, quer seja verdadeira ou falsa; mas se a é a tal coisa (em que “ a ” é um nome), as palavras “ a existe” são desprovidas de sentido. Só podemos afirmar significativamente a existência de descrições — definidas ou indefinidas; pois se “ a ” for um nome, *deve* nomear alguma coisa: o que não nomeia coisa alguma não é um nome e, portanto, se pretende ser um nome, é um símbolo desprovido de sentido, ao passo que uma descrição, como “o atual rei da França”, não se torna incapaz de ocorrer significativamente apenas em razão do fato de não descrever nada; a razão é que esse é um símbolo *complexo*, cujo significado é derivado do significado dos símbolos constituintes. Assim, quando perguntamos se Homero existiu, estamos usando a palavra “Homero” como uma descrição abreviada: podemos substituí-la por (digamos) “o autor da *Ilíada* e da *Odisséia*”. As mesmas considerações se aplicam a quase todos os usos do que parecem ser nomes próprios.

Quando descrições ocorrem em proposições, é necessário distinguir o que podemos chamar de ocorrências “primárias” e “secundárias”. A distinção abstrata é a que se segue. Uma descrição tem uma ocorrência “primária” quando a proposição em que ela ocorre resulta da substituição de x pela descrição em alguma função proposicional $\emptyset x$; uma descrição tem uma ocorrência “secundária” quando o resultado da substituição de x

pela descrição em ϕx dá somente *parte* da proposição envolvida. Um exemplo tornará isso mais claro. Consideremos “o atual rei da França é calvo”. Aqui “o atual rei da França” tem uma ocorrência primária, e a proposição é falsa. Toda proposição em que uma descrição que não descreve nada tem uma ocorrência primária é falso. Mas agora consideremos “o atual rei da França não é calvo”. Isso é ambíguo. Se tomássemos primeiro “ x é calvo” e depois substituíssemos x por “o atual rei da França” e então negássemos o resultado, a ocorrência de “o atual rei da França” é secundária e nossa proposição é verdadeira; mas se tomássemos “ x não é calvo” e substituíssemos x por “o atual rei da França”, então “o atual rei da França” tem uma ocorrência primária e a proposição é falsa. A confusão entre ocorrências primária e secundária é uma farta fonte de falácias no que diz respeito a descrições.

Descrições ocorrem em matemática sobretudo na forma de *funções descritivas*, isto é, “o termo que tem a relação R com y ”, ou “o R de y ”, como podemos dizer, por analogia com “o pai de y ” e expressões similares. Dizer “o pai de y é rico”, por exemplo, é dizer que a seguinte função proposicional de c : “ c é rico, e ‘ x gerou y ’ é sempre equivalente a ‘ x é c ’” é “verdadeiro às vezes”, ou seja, é verdadeiro para pelo menos um valor de c . Obviamente não pode ser verdadeiro para mais de um valor.

A teoria das descrições, brevemente esboçada nesse capítulo, é da máxima importância tanto em lógica quanto em teoria do conhecimento. Para os propósitos da matemática, porém, as partes mais filosóficas da teoria não são essenciais, e por isso foram omitidas da explicação anterior, que se limitou aos mais simples requisitos matemáticos.

Capítulo 17

Classes

Nesse capítulo nos ocuparemos de *os*: os habitantes de Londres, os filhos dos homens ricos, e assim por diante. Em outras palavras, estaremos tratando de *classes*. Vimos no Capítulo 2 que um número cardinal deve ser definido como uma classe de classes, e no Capítulo 3 que o número 1 deve ser definido como a classe de todas as classes unitárias, por exemplo, todas que têm apenas um membro, como diríamos não fosse o círculo vicioso. É claro que, quando o número 1 é definido como a classe de todas as classes unitárias, “classes unitárias” devem ser definidas de modo que não se suponha que saibamos o que se deve entender por “um”; de fato, elas são definidas de uma maneira estreitamente análoga à que foi usada para descrições, a saber: diz-se que uma classe x é uma classe “unitária” se a função proposicional “ x é um x ” é sempre equivalente a “ x é c ” (considerada uma função de c) não for sempre falsa, isto é, em linguagem mais comum, se houver um termo c tal que x seja um membro de x quando x for c , mas não em outros casos. Isso nos dá uma definição de uma classe unitária se já soubermos o que é uma classe em geral. Até agora, ao lidar com a aritmética, temos tratado “classe” como uma idéia primitiva. Mas, pelas razões expressas no Capítulo 13, senão por outras, não podemos aceitar “classe” como uma idéia primitiva. Devemos procurar uma definição nas mesmas linhas que a definição de descrições, isto é, uma definição que atribua um significado a proposições em cuja expressão verbal ou simbólica palavras ou símbolos que aparentemente representem classes ocorram, mas que atribuirão um sentido que elimina por completo toda menção a classe de uma análise correta de tais proposições. Poderemos dizer então que os símbolos para classes são meras conveniências, não representando objetos chamados “classes”, e que as classes são de fato, como as descrições, ficções lógicas, ou (como dizemos) “símbolos incompletos”.

A teoria das classes é menos completa que a teoria das descrições, e há razões (de que exporemos um esboço) para considerar não finalmente satisfatória a definição de classes que será sugerida. Alguma sutileza adicional parece ser necessária; mas são indiscutíveis as razões para considerar a definição que será oferecida como aproximadamente correta e na linha certa.

A primeira coisa é compreender por que classes não podem ser consideradas parte da mobília fundamental do mundo. É difícil explicar precisamente o que queremos dizer com essa afirmação, mas uma consequência que ela implica pode ser usada para elucidar seu significado. Se tivéssemos uma linguagem completamente simbólica, com uma definição para tudo o que fosse definível, os símbolos indefinidos nessa linguagem representariam simbolicamente o que quero dizer por “a mobília fundamental do mundo”. Estou sustentando que nenhum símbolo, seja para “classe” em geral ou para classes particulares, seria incluído nesse aparato de símbolos indefinidos. Por outro lado, todas as coisas particulares que há no mundo teriam de ter nomes que seriam incluídos entre os símbolos indefinidos. Poderíamos tentar evitar essa conclusão mediante o uso de descrições. Tomemos (digamos) “a última coisa que César viu antes de morrer”. Isso é uma descrição de algo particular; poderíamos usá-la (num sentido perfeitamente legítimo) com uma *definição* desse particular. Mas se “*a*” for um *nome* para o mesmo particular, uma proposição em que “*a*” ocorra não é (como vimos no capítulo anterior) idêntica ao que essa proposição se torna quando substituímos “*a*” *por* “a última coisa que César viu antes de morrer”. Se nossa linguagem não contiver o nome “*a*”, ou algum outro nome para o mesmo particular, não teremos nenhum meio de expressar a proposição que expressamos por meio de “*a*” em contraposição àquela que expressamos por meio da descrição. Assim, descrições não permitiriam a uma linguagem perfeita prescindir de nomes para todos os particulares. A esse respeito, estamos sustentando, classes diferem de particulares, e não precisam ser representadas por símbolos indefinidos. Nossa primeira tarefa é dar as razões para essa opinião.

Já vimos que classes não podem ser consideradas espécies de indivíduos, em razão da contradição acerca de classes que não são membros de si mesmas (explicada no Capítulo 13), e porque podemos provar que o número de classes é maior do que o número de indivíduos.

Não podemos tomar classes da maneira extensional *pura* como meros montes ou conglomerados. Se fôssemos tentados a isso, nos pareceria impossível compreender como pode haver uma classe como a classe nula, que não tem absolutamente nenhum membro e não pode ser considerada um “monte”; nos pareceria também muito difícil compreender como é possível que uma classe que tem apenas um membro não seja idêntica a esse membro único. Não quero afirmar, ou negar, que existam tais entidades como “montes”. Como um lógico matemático, não me compete ter uma opinião a esse respeito. Estou sustentando apenas que, caso existam tais coisas como montes, não podemos identificá-las com as classes compostas por seus constituintes.

Chegaremos muito mais perto de uma teoria satisfatória se tentarmos identificar classes com funções proposicionais. Toda classe, como explicamos no Capítulo 2, é definida por uma função proposicional verdadeira com relação aos membros da classe e falsa com relação a outras coisas. Mas se uma classe pode ser definida por uma função proposicional, pode ser definida igualmente bem por qualquer outra que seja verdadeira sempre que a primeira for verdadeira e falsa quando a primeira for falsa. Por essa razão a classe não pode ser identificada com nenhuma dessas funções proposicionais mais do que com qualquer outra — e dada uma função proposicional, haverá sempre muitas outras que serão verdadeiras quando ela for verdadeira e falsas quando ela for falsa. Podemos dizer que duas funções proposicionais são “formalmente equivalentes” quando isso acontece. Duas *proposições* são “equivalentes” quando ambas são verdadeiras ou ambas são falsas; duas funções proposicionais ϕx , ψx são “formalmente equivalentes” quando ϕx é sempre equivalente a ψx . É o fato de haver outras funções formalmente equivalentes a uma dada função que torna impossível identificar uma classe com uma função; pois queremos que as classes sejam tais que nunca duas classes distintas possam ter exatamente os mesmos membros, e portanto duas funções formalmente equivalentes tenham de determinar a mesma classe.

Quando decidimos que classes não podem ser coisas da mesma espécie que seus membros, que não podem ser meros montes ou agregados, e também que não podem ser identificadas com funções proposicionais, torna-se muito difícil ver o que elas podem ser, se devem ser mais do que ficções simbólicas. E se pudermos encontrar alguma maneira de lidar com

elas como ficções simbólicas, aumentaremos a segurança lógica de nossa posição, uma vez que evitaremos a necessidade de supor que há classes sem sermos compelidos a fazer a suposição oposta de que não há classes. Simplesmente nos absteremos de ambas as suposições. Esse é um exemplo da navalha de Occam, a saber, “entidades não devem ser multiplicadas sem necessidade”. Quando nos recusamos a afirmar que há classes, não se deve supor que estamos afirmando dogmaticamente que não há classe alguma. Somos meramente agnósticos em relação a elas: como Laplace, podemos dizer, “*je n’ai pas besoin de cette hypothèse*”.

Formulemos as condições que um símbolo deve preencher para poder servir como uma classe. Penso que as condições a seguir serão consideradas necessárias e suficientes:

(1) Toda função proposicional deve determinar uma classe, consistindo naqueles argumentos para os quais a função é verdadeira. Dada qualquer proposição (verdadeira ou falsa), digamos acerca de Sócrates, podemos imaginar Sócrates substituído por Platão ou Aristóteles ou um gorila ou um marciano ou qualquer indivíduo no mundo. Em geral, algumas dessas substituições darão uma proposição verdadeira, e algumas, uma falsa. A classe determinada consistirá em todas aquelas substituições que dão uma proposição verdadeira. Ainda temos de decidir, é claro, o que entendemos por “todas aquelas que etc.” O que estamos observando neste momento é apenas que uma classe é tornada determinada por uma função proposicional, e que toda função proposicional determina uma classe apropriada.

(2) Duas funções proposicionais formalmente equivalentes devem determinar a mesma classe, e duas não formalmente equivalentes devem determinar classes diferentes. Isto é, uma classe é determinada pelo conjunto de seus membros, e duas classes nunca podem ter o mesmo conjunto de membros. (Se uma classe é determinada por uma função ϕx , dizemos que a é um “membro” da classe se ϕa for verdadeiro.)

(3) Devemos encontrar alguma maneira de definir não só classes, mas classes de classes. Vimos no Capítulo 2 que números cardinais devem ser definidos como classes de classes. A expressão comum da matemática elementar “A combinação de n coisas m num momento” representa uma classe de classes, a saber, a classe de todas as classes de m termos que podem ser selecionadas de uma dada classe de n termos. Sem algum

método simbólico para lidar com classes de classes, a lógica matemática se desmantelaria.

(4) Sob todas as circunstâncias, deve ser sem sentido (ou falso) supor que uma classe é membro de si mesma ou não é membro de si mesma. Isso resulta da contradição que discutimos no Capítulo 8.

(5) Por fim — e esta é a condição cujo preenchimento é mais difícil —, deve ser possível fazer proposições sobre *todas* as classes compostas de indivíduos, ou sobre *todas* as classes compostas de objetos de algum “tipo” lógico. Se esse não fosse o caso, muitos usos de classes se extraviariam — por exemplo, a indução matemática. Ao definir a posteridade de um termo dado, precisamos ser capazes de dizer que um membro da posteridade pertence a *todas* as classes hereditárias a que o dado termo pertence, e isso requer o tipo de totalidade que está em questão. A razão para que haja uma dificuldade com relação a essa condição é que pode ser provado que é impossível falar de *todas* as funções proposicionais que podem ter argumentos de um dado tipo.

Para começar, vamos ignorar essa última condição e os problemas que suscita. As duas primeiras condições podem ser tomadas juntas. Elas afirmam que deve haver uma classe, nem mais e nem menos, para cada grupo de funções proposicionais formalmente equivalentes; por exemplo, a classe dos homens deve ser a mesma dos bípedes implumes ou dos animais racionais, ou das bestas humanas, ou qualquer outra característica que se possa preferir para definir ser humano. Ora, quando dizemos que duas funções proposicionais formalmente equivalentes podem não ser idênticas, embora definam a mesma classe, podemos provar a verdade dessa asserção assinalando que uma afirmação pode ser verdadeira em relação a uma função e falsa em relação a outra; por exemplo, “acredito que todos os homens são mortais” pode ser verdadeiro, ao passo que “acredito que todos os animais racionais são mortais” pode ser falso, já que posso acreditar erroneamente que a fênix é um animal racional imortal. Assim, somos levados a considerar *afirmações sobre funções*, ou (mais corretamente) *funções de funções*.

Algumas das coisas que podem ser ditas sobre uma função podem ser consideradas ditas acerca da classe definida pela função, ao passo que outras não. A afirmação “todos os homens são mortais” envolve as funções “ x é humano” e “ x é mortal”; ora, se quisermos, podemos dizer que isso envolve as classes *homens* e *mortais*. Podemos interpretar a

afirmação de ambos os modos, porque seu valor de verdade fica inalterado se substituirmos “ x é humano” ou “ x é mortal” por qualquer função formalmente equivalente. Mas, como acabamos de ver, a afirmação “acredito que todos os homens são mortais” não pode ser considerada dizendo respeito à classe determinada por nenhuma das duas funções, porque seu valor de verdade pode ser alterado pela substituição por uma função formalmente equivalente (que deixa a classe inalterada). Chamaremos uma afirmação que envolva uma função $\varnothing x$ uma função “extensional” da função $\varnothing x$, se ela for como “todos os homens são mortais”, isto é, se seu valor de verdade ficar inalterado pela substituição por qualquer função formalmente equivalente; e quando uma função de uma função não for extensional, nós a chamaremos “intensional”, de modo que “acredito que todos os homens são mortais” é uma função intensional de “ x é humano” ou “ x é mortal”. Assim, funções *extensionais* de uma função x podem, para propósitos práticos, ser consideradas funções da classe determinada por x , ao passo que funções *intensionais* não podem ser assim consideradas.

Convém observar que todas as funções *específicas* de funções que tivemos oportunidade de introduzir na lógica matemática são extensionais. Assim, por exemplo, as duas funções fundamentais das funções são: “ $\varnothing x$ é sempre verdadeiro” e “ $\varnothing x$ é verdadeiro às vezes”. Cada uma dessas tem seu valor de verdade inalterado se $\varnothing x$ for substituído por qualquer função formalmente equivalente. Na linguagem das classes, se α for a classe determinada por $\varnothing x$, “ $\varnothing x$ é sempre verdadeiro” é equivalente a “tudo é membro de α ”, e “ $\varnothing x$ é verdadeiro às vezes” é equivalente a “ α tem membros” ou (melhor) “ α tem pelo menos um membro”. Tomemos novamente a condição, com que lidamos no capítulo anterior, para a existência de “o termo que satisfaz $\varnothing x$ ”. A condição é que haja um termo c tal que $\varnothing x$ seja sempre equivalente a “ x é c ”. Isso é obviamente extensional. É equivalente à asserção de que a classe definida pela função $\varnothing x$ é uma classe unitária, isto é, uma classe que tem 1 membro; em outras palavras, uma classe que é membro de 1.

Dada uma função que pode ou não ser extensional, podemos sempre derivar dela uma função conexa e certamente extensional da mesma função, mediante o seguinte plano: suponhamos que nossa função original de uma função seja uma que atribui a $\varnothing x$ a propriedade f ; depois consideremos a asserção “há uma função que tem a propriedade f e é

formalmente equivalente a ϕx ". Essa é uma função extensional de ϕx ; é verdadeira quando nossa afirmação original é verdadeira, e é formalmente equivalente à função original de ϕx se essa função original for extensional; mas quando a função original é intensional, a nova é verdadeira com mais frequência que a antiga. Por exemplo, considere novamente "acredito que todos os homens são mortais", visto como uma função de " x é humano". A função extensional derivada é: "há uma função formalmente equivalente a ' x é humano' e tal que eu acredito que tudo o que a satisfaz é mortal." Isso continua verdadeiro quando substituímos " x é humano" por " x é um animal racional", mesmo que eu acredite, falsamente, que a fênix é racional e imortal.

Damos o nome "função extensional derivada" à função construída como anteriormente, a saber, à função: "há uma função que tem a propriedade f e é formalmente equivalente a ϕx ", em que a função original era "a função ϕx tem a propriedade f ".

Podemos considerar que a função extensional derivada tem por seu argumento a classe determinada pela função ϕx , e afirma f sobre essa classe. Isso pode ser tomado como a definição de uma proposição sobre uma classe, isto é, podemos definir: Afirmar que "a classe determinada pela função ϕx tem a propriedade f " é afirmar que ϕx satisfaz a função extensional derivada de f .

Isso dá um sentido a qualquer afirmação sobre uma classe que possa ser feita de maneira significativa acerca de uma função; e veremos que, tecnicamente, gera o resultado necessário para tornar uma teoria simbolicamente satisfatória.¹

O que acabamos de dizer com relação à definição de classes é suficiente para satisfazer nossas quatro primeiras condições. A maneira pela qual ela assegura a terceira e a quarta, a saber, a possibilidade de classes de classes e a impossibilidade de uma classe ser ou não ser membro de si mesma, é um tanto técnica; é explicada em *Principia Mathematica*, mas pode ser dada como certa aqui. O resultado é que, exceto para nossa quinta condição, podemos considerar nossa tarefa concluída. Mas essa condição — ao mesmo tempo a mais importante e a mais difícil — não é preenchida em virtude de coisa alguma que já tenhamos dito. A dificuldade está ligada à teoria dos tipos, e deve ser brevemente discutida.²

Vimos no Capítulo 13 que há uma hierarquia de tipos lógicos, e que é uma falácia permitir que um objeto pertencente a um desses tipos

substitua um objeto pertencente a outro. Ora, não é difícil mostrar que as várias funções que podem tomar um objeto dado a como argumento não são todas de um único tipo. Podemos chamá-las todas funções de a . Podemos tomar primeiro aquelas entre elas que não envolvem referência a nenhuma coleção de funções; vamos chamá-las de “funções de a predicativas”. Se avançarmos agora para funções que envolvem referência à totalidade das funções de a predicativas, incorreremos numa falácia se as considerarmos sendo do mesmo tipo que as funções de a predicativas. Tomemos uma afirmação banal como “ a é um francês típico”. Como definiremos um “francês típico”? Podemos defini-lo como um “que possua todas as qualidades possuídas pela maioria dos franceses”. No entanto, a menos que limitemos “todas as qualidades” àquelas que não envolvem uma referência a uma totalidade de qualidades, teremos de observar que os franceses em sua maioria não são “típicos” no sentido apresentado, e portanto a definição mostra que não ser típico é essencial para um francês típico. Isso não é uma contradição lógica, uma vez que não há nenhuma razão para que haja algum francês típico; mas ilustra a necessidade de separar qualidades que envolvem referência a uma totalidade de qualidades daquelas que não o fazem.

Sempre que, por afirmações acerca de “todos” ou “alguns” dos valores que uma variável pode assumir de maneira significativa, geramos um novo objeto, esse novo objeto não deve estar entre os valores que nossa variável anterior podia assumir, uma vez que, se estivesse, a totalidade de valores entre os quais a variável poderia variar só seria definível em termos de si própria, e ficaríamos envolvidos num círculo vicioso. Por exemplo, se digo “Napoleão tinha todas as qualidades que fazem um grande general”, devo definir “qualidades” de tal maneira que isso não inclua o que estou dizendo agora, isto é, “ter todas as qualidades que fazem um grande general” não deve ser em si mesmo uma qualidade no sentido suposto. Isso é bastante óbvio, e é o princípio que leva à teoria dos tipos, pela qual paradoxos de círculo vicioso são evitados. Quando aplicadas a funções de a , podemos supor que “qualidades” deve significar “funções predicativas”. Assim, quando digo “Napoleão tinha todas as qualidades etc.”, quero dizer “Napoleão satisfazia todas as funções predicativas etc.”. Essa afirmação atribui uma propriedade a Napoleão, mas não uma propriedade predicativa; dessa maneira, escapamos ao círculo vicioso. Mas sempre que “todas as funções que” ocorre, a função em questão deve ser limitada a um

único tipo para que o círculo vicioso seja evitado; e, como Napoleão e o francês típico mostraram, o tipo não é tornado determinado por aquele do argumento. Precisaríamos de uma discussão muito mais extensa para expor essa idéia por completo, mas o que foi dito pode bastar para deixar claro que as funções que podem tomar um dado argumento são de uma série infinita de tipos. Poderíamos, mediante vários estratagemas lógicos, construir uma variável que percorresse o primeiro n desses tipos, em que n é finito, mas não podemos construir uma variável que percorresse todos eles, e, se pudéssemos, esse simples fato geraria de imediato um novo tipo de função com os mesmos argumentos, e desencadearia de novo o mesmo processo.

Chamamos funções de a predicativas o *primeiro* tipo de funções de a ; chamamos as funções de a que envolvem referência à totalidade do primeiro tipo de *segundo* tipo; e assim por diante. Nenhuma função de a variável pode percorrer todos esses diferentes tipos: deve se limitar a algum tipo definido.

Essas considerações são relevantes para nossa definição da função derivada extensional. Falamos nesse caso de uma “função formalmente equivalente a $\emptyset x$ ”. É necessário decidir quanto ao tipo de nossa função. Qualquer decisão serve, mas algumas são inevitáveis. Chamemos a suposta função formalmente equivalente ψ . Portanto, ψ aparece como uma variável, e deve ser de algum tipo determinado. O que sabemos necessariamente sobre o tipo de $\emptyset$ é apenas que toma argumentos de um dado tipo — é (digamos) uma função de a . Mas isso, como acabamos de ver, não determina seu tipo. Para sermos capazes (como nosso quinto requisito exige) de lidar com *todas* as classes cujos membros são do mesmo tipo que a , devemos ser capazes de definir todas essas classes por meio de funções de um único tipo; isto é, deve haver algum tipo de função de a , digamos o *enésimo*, tal que qualquer função de a seja formalmente equivalente a alguma função de a do *enésimo* tipo. Se esse for o caso, então qualquer função extensional que se aplique a todas as funções de a do *enésimo* tipo se aplicará a qualquer função de a . É principalmente como um meio técnico de materializar uma suposição que conduza a esse resultado que as classes são úteis. A suposição é chamada o “axioma da redutibilidade” e pode ser expressa da seguinte maneira: “Há um tipo (digamos r) de funções de a que, dada qualquer função de a , é formalmente equivalente a alguma função do tipo em questão.”

Se esse axioma for admitido, usaremos funções desse tipo ao definir nossa função extensional associada. Afirmarções sobre todas as classes a (isto é, todas as classes definidas por funções de a) podem ser reduzidas a afirmações sobre todas as funções de a do tipo r . Enquanto apenas funções extensionais de funções estiverem envolvidas, isso nos dá na prática resultados que de outro modo teriam exigido a noção impossível de “todas as funções de a ”. Uma região particular em que isso é vital é a indução matemática.

O axioma da redutibilidade envolve tudo o que é realmente essencial na teoria das classes. Vale a pena, portanto, perguntar se há alguma razão para supor que ele é verdadeiro.

Esse axioma, como o axioma multiplicativo e o axioma da infinidade, é necessário para certos resultados, mas não para a simples existência do raciocínio dedutivo. A teoria da dedução, como explicada no Capítulo 14, e as leis para proposições que envolvem “todos” e “alguns” são a própria textura do raciocínio matemático: sem elas, ou algo semelhante a elas, não só não obteríamos os mesmos resultados, como não obteríamos resultado algum. Não podemos usá-las como hipóteses, e deduzir conseqüências hipotéticas, porque elas são tanto regras de dedução quanto premissas. Devem ser absolutamente verdadeiras, do contrário o que deduzimos de conformidade com elas sequer se segue das premissas. Por outro lado, o axioma da redutibilidade, como nossos dois axiomas matemáticos anteriores, poderia perfeitamente ser formulado como uma hipótese sempre que é usado, em vez de ser admitido como realmente verdadeiro. Podemos deduzir suas conseqüências hipoteticamente; podemos também deduzir as conseqüências de supô-lo falso. Ele é, portanto, apenas conveniente, não necessário. E, em vista da complicação da teoria dos tipos e da incerteza de tudo exceto seus princípios mais gerais, por enquanto é impossível dizer se não poderia haver alguma maneira de dispensar por completo o axioma da redutibilidade. No entanto, supondo a correção da teoria esboçada anteriormente, que podemos dizer quanto à verdade ou falsidade do axioma?

O axioma, podemos observar, é uma forma generalizada da identidade dos indiscerníveis de Leibniz. Esse supunha, como um princípio lógico, que dois sujeitos diferentes devem diferir quanto a seus predicados. Ora, predicados são somente algumas entre o que chamamos “funções predicativas”, que incluirão também relações com termos dados, e várias

propriedades que não devem ser consideradas predicados. Assim a suposição de Leibniz é muito mais estrita e estreita que a nossa. (Não, é claro, de acordo com a lógica *dele*, que considerava *todas* as proposições redutíveis à forma sujeito-predicado.) Mas, até onde posso ver, não há nenhuma boa razão para se acreditar na forma de Leibniz. Seria perfeitamente possível, como uma possibilidade lógica abstrata, haver duas coisas que tivessem exatamente os mesmos predicados no sentido estrito em que temos usado a palavra “predicado”. Como fica nosso axioma quando passamos além dos predicados no sentido estreito? No mundo real, parece não haver como duvidar de sua verdade empírica no tocante a particulares, em virtude da diferenciação espaço-temporal: dois particulares nunca têm exatamente as mesmas relações espaciais e temporais com todos os outros particulares. Mas isso é, por assim dizer, um acidente, um fato sobre o mundo em que por acaso nos encontramos. A lógica pura, e a matemática pura (que é a mesma coisa), pretendem ser verdadeiras, para usar a fraseologia de Leibniz, em todos os mundos possíveis, não apenas neste mundo confuso e acidentado em que o acaso nos aprisionou. Há certa altivez que o lógico deve preservar: não deve condescender em derivar argumentos das coisas que vê à sua volta.

Encarando-o desse ponto de vista estritamente lógico, não me parece haver nenhuma razão para acreditar que o axioma da redutibilidade é logicamente necessário, o que seria o sentido de dizer que é verdadeiro em todos os mundos possíveis. A admissão desse axioma num sistema de lógica é, portanto, um defeito, ainda que ele seja empiricamente verdadeiro. É por essa razão que a teoria das classes não pode ser considerada tão completa quanto a teoria das descrições. Há necessidade de mais trabalho sobre a teoria dos tipos, na esperança de se chegar a uma doutrina das classes que não exija essa suposição dúbia. Mas é razoável considerar que a teoria esboçada no presente capítulo é correta em suas linhas principais, isto é, em sua redução de proposições nominalmente acerca de classe a proposições acerca de suas funções definidoras. Evitar classes como entidades por meio desse método deve, ao que parece, ser válido em princípio, embora os detalhes possam ainda requerer ajustes. É porque isso parece indubitável que incluímos aqui a teoria das classes, apesar de nosso desejo de excluir, tanto quanto possível, tudo o que parecesse aberto a dúvidas sérias.

A teoria das classes, como esboçada anteriormente, reduz-se a um axioma e uma definição. No interesse da clareza, vamos repeti-los. O axioma é:

Há um tipo τ tal que se $\emptyset$ for uma função que pode tomar um dado objeto como argumento, então há uma função ψ do tipo τ que é formalmente equivalente a $\emptyset$.

A definição é:

Se $\emptyset$ for uma função que pode tomar um dado objeto como argumento, e τ o tipo mencionado no axioma apresentado, então dizer que a classe determinada por $\emptyset$ tem a propriedade f é dizer que há uma função do tipo τ , que é formalmente equivalente a $\emptyset$ e tem a propriedade f .

Capítulo 18

Matemática e lógica

Historicamente falando, a matemática e a lógica foram estudos inteiramente distintos. A matemática esteve ligada à ciência; a lógica, ao grego. Ambas, porém, desenvolveram-se nos tempos modernos: a lógica tornou-se mais matemática e a matemática tornou-se mais lógica. A consequência é que agora se tornou inteiramente impossível traçar uma linha entre as duas; de fato, as duas são uma. Diferem como um menino e um homem: a lógica é a juventude da matemática, e a matemática, a maturidade da lógica. Essa idéia deixa indignados os lógicos que, tendo passado seu tempo estudando dos textos clássicos, são incapazes de acompanhar um raciocínio simbólico, e os matemáticos que aprenderam uma técnica sem se dar ao trabalho de indagar sobre seu significado ou justificação. Hoje em dia, felizmente, os dois tipos estão ficando mais raros. Uma parte tão grande do trabalho matemático moderno situa-se na fronteira da lógica, uma parte tão grande da lógica é simbólica e formal, que a relação muito estreita entre a lógica e a matemática tornou-se óbvia para todo estudante instruído. A prova da identidade de ambas, claro, é uma questão de detalhe: começando com premissas que seriam universalmente admitidas como pertencentes à lógica, e chegando, por dedução, a resultados que pertencem de maneira igualmente óbvia à matemática, descobrimos que não há ponto algum em que uma linha nítida possa ser traçada, com a lógica à esquerda e a matemática à direita. Se ainda houver quem não admite a identidade entre a lógica e a matemática, podemos desafiá-los a indicar em que ponto, nas sucessivas definições e deduções de *Principia Mathematica*, eles consideram que termina a lógica começa a matemática. Ficará evidente então que qualquer resposta deverá ser completamente arbitrária.

Nos capítulos anteriores, começando pelos números naturais, definimos primeiro “número cardinal” e mostramos como generalizar a concepção de

número; depois analisamos as concepções envolvidas na definição, até que tratamos dos fundamentos da lógica. Num tratamento sintético, dedutivo, esses fundamentos vêm em primeiro lugar, e os números naturais só são alcançados após uma longa jornada. Esse tratamento, embora formalmente mais correto que aquele que adotamos, é mais difícil para o leitor, porque os conceitos lógicos e as proposições fundamentais com que começa são distantes e estranhos se comparados com os números naturais. Além disso, eles representam a atual fronteira do conhecimento, além da qual situa-se o ainda desconhecido; e o domínio do conhecimento sobre eles ainda não é muito seguro.

Costumava-se dizer que a matemática é a ciência da “quantidade”. “Quantidade” é um termo vago, mas no interesse da argumentação podemos substituí-lo pela palavra “número”. A afirmação de que a matemática é a ciência do número seria falsa de duas diferentes maneiras. Por um lado, há ramos reconhecidos da matemática que nada têm a ver com número — toda a geometria que não usa coordenadas ou medidas, por exemplo: a geometria projetiva e descritiva, até o ponto em que coordenadas são introduzidas, nada tem a ver com número, ou sequer com quantidade no sentido de *maior* e *menor*. Por outro lado, por meio da definição de cardinais, da teoria da indução e das relações ancestrais, da teoria geral das séries e da definição das operações aritméticas, tornou-se possível generalizar grande parte do que costumava ser provado apenas em conexão com números. O resultado é que o que era outrora o único estudo da aritmética dividiu-se agora em vários estudos separados, dos quais nenhum diz respeito especialmente a números. As propriedades mais elementares dos números dizem respeito a relações um-um e a similaridade entre classes. A adição diz respeito à construção de classes mutuamente exclusivas respectivamente similares a um conjunto de classes que não sabemos serem mutuamente exclusivas. A multiplicação é incorporada na teoria das “seleções”, isto é, de certo tipo de relações um-muitos que gera toda a teoria da indução matemática. As propriedades ordinais dos vários tipos de séries de números, bem como os elementos da teoria da continuidade das funções e dos limites das funções podem ser generalizados de modo que não mais envolva qualquer referência essencial a números. É um princípio, em todo raciocínio formal, generalizar ao máximo, uma vez que com isso podemos assegurar que um dado processo de dedução tenha resultados mais amplamente aplicáveis; portanto, ao

generalizar assim o raciocínio da aritmética, estamos somente seguindo um preceito universalmente admitido na matemática. E ao generalizar assim, criamos, de fato, um conjunto de novos sistemas dedutivos, em que a aritmética tradicional é ao mesmo tempo dissolvida e ampliada; mas se devemos dizer que algum desses novos sistemas dedutivos — por exemplo, a teoria das seleções — pertence à lógica ou à matemática é algo inteiramente arbitrário e não passível de ser decidido racionalmente.

Somos, portanto, levados a encarar de frente a questão: que matéria é essa que pode ser chamada indiferentemente matemática ou lógica? Há algum meio que nos permita definir isso?

Determinadas características dessa matéria são claras. Para começar, nessa matéria não lidamos com coisas particulares ou propriedades particulares: lidamos formalmente com o que pode ser dito acerca de *qualquer* coisa ou *qualquer* propriedade. Estamos prontos a dizer que um mais um são dois, porém não que Sócrates mais Platão são dois, porque, em nossa condição de lógicos ou de matemáticos puros, nunca ouvimos falar de Sócrates ou Platão. Um mundo em que tais indivíduos nunca tenham existido ainda seria um mundo em que um mais um seriam dois. Não nos é dado, em nossa condição de matemáticos puros ou lógicos, mencionar coisa alguma, porque se o fizemos estaremos introduzindo algo irrelevante e não formal. Podemos tornar isso claro aplicando-o ao caso do silogismo. A lógica tradicional diz: “Todos os homens são mortais, Sócrates é um homem, portanto Sócrates é mortal.” Fica claro agora que o que *pretendemos* afirmar, para começar, é apenas que as premissas implicam a conclusão, não que as premissas e a conclusão sejam realmente verdadeiras; mesmo o lógico mais tradicional assinala que a verdade efetiva das premissas é irrelevante para a lógica. Assim, a primeira mudança a ser feita no silogismo tradicional apresentada é expressá-lo na forma: “Se todos os homens são mortais e Sócrates é um homem, então Sócrates é mortal.” Podemos observar agora que isso pretende comunicar que esse raciocínio é válido em virtude de sua *forma*, não em virtude dos termos particulares que nela ocorrem. Se tivéssemos omitido “Sócrates é um homem” de nossas premissas, teríamos um raciocínio não formal, admissível somente porque Sócrates é de fato um homem; nesse caso não teríamos podido generalizar esse raciocínio. Mas quando, como no caso citado, o raciocínio é *formal*, nada depende dos termos que nele ocorrem. Dessa maneira, podemos substituir *homens* por

α , mortais por β e Sócrates por x , em que α e β são quaisquer classes e x é qualquer indivíduo. Chegamos, então, à afirmação: “Quaisquer que sejam os valores de x , α e β , se todos os α 's são β 's e x é um α , então x é um β ”; em outras palavras, “a função proposicional ‘se todos os α 's são β 's e x é um α , então x é um β ’ é sempre verdade”. Aqui finalmente temos uma proposição da lógica — aquela que é *sugerida* pela afirmação tradicional acerca de Sócrates, homens e mortais.

É claro que, se nosso objetivo for chegar ao raciocínio *formal*, chegaremos sempre em última instância a afirmações como a apresentada, em que nenhuma coisa ou propriedade real é mencionada; isso acontecerá por meio de nosso mero desejo de não perder nosso tempo provando num caso particular o que pode ser provado em geral. Seria ridículo desenvolver um longo raciocínio acerca de Sócrates, e depois desenvolver precisamente o mesmo de novo acerca de Platão. Se nosso raciocínio é tal que (digamos) se aplica a todos os homens, vamos prová-lo com relação a x , com a hipótese “se x é um homem”. Com essa hipótese, o raciocínio conservará sua validade hipotética mesmo quando x não for um homem. Mas agora verificaremos que nosso raciocínio ainda seria válido se, em vez de supor que x é um homem, supuséssemos que x é um macaco, um ganso ou um primeiro-ministro. Não perderemos, portanto, nosso tempo tomando como premissa “ x é um homem”, mas tomaremos “ x é um α ”, em que α é qualquer classe de indivíduos, ou ϕx , em que ϕ é qualquer função proposicional de algum tipo designado. Assim, a ausência de toda menção a coisas ou propriedades singulares na lógica ou na matemática pura é um resultado necessário do fato de seu estudo ser, como dizemos, “puramente formal”.

Neste ponto defrontamo-nos com um problema mais fácil de formular que de resolver. O problema é: “Quais são os constituintes de uma proposição lógica?” Não sei a resposta, mas proponho-me a explicar como o problema surge.

Tomemos (digamos) a proposição “Sócrates existiu antes de Aristóteles”. Aqui parece óbvio que temos uma relação entre dois termos, e que os constituintes da proposição (bem como do fato correspondente) são simplesmente os dois termos e a relação, isto é, Sócrates, Aristóteles e *antes*. (Ignoro o fato de que Sócrates e Aristóteles não são simples; também o fato de que o que aparentemente são seus nomes, são na realidade descrições truncadas. Nenhum desses fatos é relevante para

nossa presente questão.) Podemos representar a forma geral de tais proposições por “ $x R y$ ”, que pode ser lido como “ x tem a relação R com y ”. Essa forma geral pode ocorrer em proposições lógicas, mas não algum caso particular dela. Devemos então inferir que a forma geral é ela própria um constituinte de tais proposições lógicas?

Dada uma proposição como “Sócrates existiu antes Aristóteles”, temos certos constituintes e também certa forma. Mas a forma não é ela mesma um novo constituinte; se fosse, precisaríamos de uma nova forma para abarcar tanto a ela quanto aos outros constituintes. Podemos, de fato, transformar *todos* os constituintes de uma proposição em variáveis, mantendo ao mesmo tempo a forma inalterada. Isso é o que fazemos quando usamos um esquema como “ $x R y$ ”, que representa qualquer uma de certa classe de proposições, a saber, aquela que afirma relações entre dois termos. Podemos passar a asserções gerais, como “ $x R y$ é verdadeiro às vezes” — isto é, há casos em que existem relações duais. Essa asserção pertencerá à lógica (ou à matemática) no sentido em que estamos usando a palavra. Mas nela não mencionamos nenhuma coisa ou relação particular; nenhuma coisa ou relação particular pode jamais fazer parte de uma proposição de lógica pura. Ficamos reduzidos a *formas* puras como os únicos constituintes possíveis das proposições lógicas.

Não desejo afirmar positivamente que formas puras — por exemplo, a forma “ $x R y$ ” — fazem realmente parte de proposições do tipo que estamos considerando. A questão da análise de tais proposições é difícil, com considerações conflitantes de um lado e de outro. Não podemos nos envolver nessa questão agora, mas podemos aceitar, como uma primeira aproximação, a idéia de que *formas* são o que faz parte das proposições lógicas como seus constituintes. E podemos explicar (embora não definir formalmente) o que queremos dizer por “forma” de uma proposição da seguinte maneira:

A “forma” de uma proposição é o que, nela, permanece inalterado quando todos os constituintes da proposição são substituídos por outros.

Assim, “Sócrates é anterior a Aristóteles” tem a mesma forma que “Napoleão é maior do que Wellington”, embora todos os constituintes das duas proposições sejam diferentes.

Podemos então estabelecer, como uma característica necessária (embora não suficiente) das proposições lógicas ou matemáticas, que elas devem ser tais que possam ser obtidas de uma proposição que não contenha

nenhuma variável (isto é, nenhuma palavra como *todos*, *alguns*, *um*, *o* etc.) pela transformação de cada constituinte numa variável e afirmando-se que o resultado é sempre verdadeiro ou às vezes verdadeiro, ou que é sempre verdadeiro com respeito a algumas das variáveis, ou que é verdadeiro às vezes com respeito a outras, ou qualquer variante dessas formas. Uma outra maneira de expressar a mesma coisa é dizer que a lógica (ou a matemática) diz respeito unicamente a *formas* e diz respeito a elas unicamente na maneira de afirmar que são sempre ou às vezes verdadeiras — com todas as permutações de “sempre” e “às vezes” que possam ocorrer.

Há em todas as línguas algumas palavras cuja única função é indicar forma. Essas palavras, geralmente, são mais comuns em línguas que têm menos inflexões. Tomemos “Sócrates é humano”. Aqui, “é” não é um constituinte da proposição, mas indica meramente a forma sujeito–predicado. De maneira similar, em “Sócrates é anterior a Aristóteles”, o “é” e o “a” indicam meramente forma; a proposição é a mesma que “Sócrates precede Aristóteles”, em que essas palavras desapareceram e a forma é indicada de outra maneira. A forma, via de regra, *pode* ser indicada de outra maneira que por palavras específicas: a ordem das palavras pode assegurar quase tudo o que se deseja. Mas não se deve enfatizar demais esse princípio. Por exemplo, é difícil ver como poderíamos expressar convenientemente formas moleculares de proposições (isto é, o que chamamos “funções de verdade”) sem absolutamente nenhuma palavra. Vimos no Capítulo 14 que uma palavra ou símbolo é suficiente para esse propósito, a saber, uma palavra ou símbolo que expressem *incompatibilidade*. Mas sem nenhum poderíamos encontrar dificuldades. Esse, no entanto, não é o ponto importante para nosso objetivo no momento. O importante para nós é observar que a forma pode ser o único interesse de uma proposição geral, mesmo quando nenhuma palavra ou símbolo nessa proposição designe a forma. Se desejamos falar acerca da forma em si mesma, devemos ter uma palavra para ela; mas se, como na matemática, desejamos falar acerca de todas as proposições que têm a forma, em geral se verificará que uma palavra para a forma não é indispensável; provavelmente em teoria ela *nunca* é indispensável.

Supondo — como penso que podemos fazer — que as formas das proposições *podem* ser representadas pelas formas das proposições em que

elas são expressas sem nenhuma palavra especial para formas, chegaríamos a uma linguagem em que tudo o que fosse formal pertenceria à sintaxe e não ao vocabulário. Em tal linguagem, poderíamos expressar *todas* as proposições da matemática, mesmo que não conhecêssemos uma única palavra da linguagem. A linguagem da lógica matemática, se ela fosse desenvolvida, seria tal linguagem. Teríamos símbolos para variáveis, como “ x ”, “ R ” e “ y ”, arranjos de vários modos; e o modo de arranjo indicaria que se estava afirmando a verdade de alguma coisa para todos os valores ou alguns valores das variáveis. Não precisaríamos conhecer nenhuma palavra, porque elas só seriam necessárias para conferir valores às variáveis, o que é a tarefa do matemático aplicado, não do matemático puro ou do lógico. Uma das marcas de uma proposição da lógica é que, dada uma linguagem adequada, uma tal proposição pode ser afirmada nessa linguagem por uma pessoa que conheça a sintaxe sem conhecer uma só palavra do vocabulário.

Mas, afinal de contas, há palavras que expressam forma, como “é” e “a”. E em todo simbolismo até hoje inventado para a lógica matemática, há símbolos que possuem significados formais constantes. Podemos tomar como exemplo o símbolo para incompatibilidade, que é empregado no estabelecimento de funções de verdade. Tais palavras ou símbolos podem ocorrer na lógica. A questão é: como devemos defini-los?

Tais palavras ou símbolos expressam as chamadas “constantes lógicas”. As constantes lógicas podem ser definidas exatamente como definimos formas; de fato, são em essência a mesma coisa. Uma constante lógica fundamental será aquela que é comum a várias proposições, qualquer das quais pode resultar da outra por substituições de um termo por outro. Por exemplo, “Napoleão é maior do que Wellington” resulta de “Sócrates é anterior a Platão”, pela substituição de “Sócrates” por “Napoleão”, de “Platão” por “Wellington” e de “anterior” por “maior”. Algumas proposições podem ser obtidas dessa maneira tomando-se por base o protótipo “Sócrates é anterior a Platão” e algumas não; as que podem são as que são da forma “ $x R y$ ”, isto é, expressam relações duais. Não podemos obter do protótipo citado, por substituição termo por termo, proposições como “Sócrates é humano” ou “os atenienses deram cicuta a Sócrates”, porque a primeira é da forma sujeito–predicado e a segunda expressa uma relação entre três termos. Se devemos ter palavras em nossa linguagem lógica pura, elas devem ser tais que expressem “constantes

lógicas”, e “constantes lógicas” serão sempre o que é comum a um grupo de proposições deriváveis umas das outras, da maneira apresentada anteriormente, por substituições termo por termo (ou derivadas do que é comum a um tal grupo). E isso que há em comum é o que chamamos “forma”.

Nesse sentido, todas as “constantes” que ocorrem na matemática pura são constantes lógicas. O número 1, por exemplo, é derivado de proposições da forma: “Há um termo c tal que $\forall x$ é verdadeiro quando, e somente quando, x é c .” Essa é uma função de $\forall$, e várias diferentes proposições resultam da atribuição de diferentes valores a $\forall$. Podemos (com uma pequena omissão de passos intermediários não relevante para nosso objetivo presente) tomar a função citada de $\forall$ como o que se quer dizer por “a classe determinada por $\forall$ é uma classe unitária” ou “a classe determinada por $\forall$ é membro de 1” (1 sendo uma classe de classes). Dessa maneira, proposições em que 1 ocorre adquirem um significado que é derivado de certa forma lógica constante. E veremos que o mesmo ocorre com todas as constantes matemáticas: todas são constantes lógicas, ou abreviações simbólicas cujo uso completo num contexto apropriado é definido por meio de constantes lógicas.

No entanto, embora todas as proposições lógicas (ou matemáticas) possam ser inteiramente expressas em termos de constantes lógicas juntamente com variáveis, não é verdade que, inversamente, todas as proposições que podem ser expressas dessa maneira sejam lógicas. Até agora encontramos um critério necessário, porém não suficiente, para proposições matemáticas. Definimos suficientemente o caráter das *idéias* primitivas em termos das quais todas as idéias da matemática podem ser *definidas*, mas não das proposições primitivas com base nas quais todas as proposições da matemática podem ser *deduzidas*. Essa é uma questão mais difícil, cuja resposta completa ainda não se conhece.

Podemos tomar o axioma da infinidade como um exemplo de proposição que, embora possa ser enunciada em termos lógicos, não pode ser afirmada como verdadeira pela lógica. Todas as proposições da lógica têm uma característica que se costumava expressar dizendo que eram analíticas, ou que seus contraditórios eram incoerentes. Esse modo de afirmação, no entanto, não é satisfatório. A lei da contradição é meramente uma entre as proposições lógicas; não tem nenhuma preeminência especial; e a prova de que a contradição de alguma

proposição é incoerente provavelmente requer outros princípios de dedução afora a lei da contradição. Apesar disso, a característica das proposições lógicas que estamos procurando é tal que aqueles que disseram que ela consistia em dedutibilidade da lei da contradição perceberam e pretenderam definir. Essa característica, que por ora podemos chamar *tautologia*, obviamente não pertence à asserção de que o número de indivíduos no universo é n , não importa que número n seja. Mas para a diversidade dos tipos, seria possível provar logicamente que há classes de n termos, em que n é qualquer número inteiro finito; ou até que há classe de n_0 termos. Mas, por causa dos tipos, essas provas, como vimos no Capítulo 13, são falaciosas. Ficamos reduzidos a observações empíricas para determinar se o número de indivíduos no mundo é n . Entre os mundos “possíveis” no sentido leibniziano, haverá mundos que tenham um, dois, três. . . indivíduos. Não parece haver nem mesmo nenhuma necessidade lógica de que haja sequer um indivíduo¹ — por que, de fato, deveria haver algum mundo? A prova ontológica da existência de Deus, se fosse válida, estabeleceria a necessidade lógica de pelo menos um indivíduo. Mas ela é geralmente reconhecida como inválida, e de fato repousa sobre uma idéia errônea de existência — isto é, não compreende que só se pode afirmar a existência de algo descrito, não de algo nomeado, de modo que não faz sentido argumentar a partir de “isso é tal coisa” e “tal coisa existe”. Se rejeitamos o argumento ontológico, somos levados, ao que parece, a concluir que a existência de um mundo é um acidente — isto é, não é logicamente necessária. Se for assim, nenhum princípio da lógica pode afirmar “existência”, exceto sob uma hipótese, ou seja, nenhum deles pode ser da forma “a função proposicional tal é verdadeira às vezes”. Proposições com essa forma, quando ocorrem na lógica, terão de ocorrer como hipóteses ou conseqüências de hipóteses, não como proposições afirmadas completas. As proposições afirmadas completas da lógica serão todas do tipo que afirma que alguma função proposicional é *sempre* verdade. Por exemplo, é sempre verdade que se p implica q e q implica r , então p implica r , ou que, se todos os α 's são β 's e x é um α , então x é um β . Tais proposições podem ocorrer na lógica, e sua verdade é independente da existência do universo. Podemos estabelecer que, se não houvesse nenhum universo, *todas* as proposições gerais seriam verdadeiras, pois o contraditório de uma proposição geral (como vimos no Capítulo 15) é uma

proposição que afirma existência, e seria portanto sempre falsa se nenhum universo existisse.

As proposições lógicas são tais que podem ser conhecidas *a priori*, sem estudo do mundo real. Só com base em um estudo de fatos empíricos sabemos que Sócrates é um homem, mas sabemos que o silogismo está correto em sua forma abstrata (isto é, quando é formulado em termos de variáveis) sem precisar fazer nenhum apelo à experiência. Isso é uma característica não das proposições lógicas em si mesmas, mas do modo pelo qual as conhecemos. Tem, contudo, uma relação com a questão de qual pode ser a natureza delas, uma vez que há alguns tipos de proposições que seria muito difícil supor que teríamos como conhecê-la sem experiência.

Está claro que devemos procurar a definição de “lógica” ou “matemática” tentando dar uma nova definição da velha noção de proposições “analíticas”. Embora não possamos mais nos satisfazer em definir proposições lógicas como aquelas que se seguem da lei da contradição, ainda podemos e devemos admitir que elas são uma classe de proposições inteiramente diferente daquelas que chegamos a conhecer empiricamente. Todas elas têm a característica que, um momento atrás, concordamos em chamar “tautologia”. Isso, combinado com o fato de que podem ser inteiramente expressas em termos de variáveis e constantes lógicas (uma constante lógica sendo algo que permanece constante numa proposição mesmo quando *todos* os seus constituintes são alterados), dará a definição de lógica ou matemática pura. Por enquanto, não sei como definir “tautologia”.² Seria fácil propor uma definição que poderia ser satisfatória por algum tempo; mas não sei de nenhuma que sinta ser satisfatória, apesar de saber muito bem que característica é preciso definir. Neste ponto, portanto, por enquanto, chegamos à fronteira do conhecimento em nossa viagem de retorno aos fundamentos lógicos da matemática.

Chegamos agora ao fim de nossa introdução um tanto sumária à filosofia matemática. É impossível transmitir adequadamente as idéias envolvidas nessa matéria enquanto nos abstermos do uso de símbolos lógicos. Como a linguagem natural não tem palavras que expressem de maneira natural exatamente o que desejamos expressar, é necessário, enquanto aderirmos à linguagem comum, forçar as palavras, atribuindo-lhes sentidos não usuais; e o leitor certamente, após algum tempo, senão

de saída, voltará aos poucos a associar os significados usuais às palavras, chegando assim a noções erradas do que se pretendeu dizer. Ademais, a gramática e a sintaxe comuns são extraordinariamente enganosas. Esse é o caso, por exemplo, com relação aos números; “dez homens” é gramaticalmente da mesma forma que “grandes homens”, de modo que se poderia pensar que dez é um adjetivo que qualifica “homens”. Esse é o caso, também, onde quer que funções proposicionais estejam envolvidas, e em particular no que diz respeito a existência e descrições. Como a linguagem é enganosa, bem como porque é difusa e inexata quando aplicada à lógica (para a qual nunca se destinou), o simbolismo lógico é absolutamente necessário para qualquer tratamento exato ou completo de nossa matéria. Espera-se, portanto, que aqueles leitores que desejem adquirir uma mestria dos princípios da matemática não se esquivarão ao trabalho de dominar os símbolos — um trabalho que é, de fato, muito menor do que se poderia pensar. Como o apressado exame feito anteriormente deve ter evidenciado, há inumeráveis problemas não resolvidos na matéria, e muito trabalho precisa ser feito. Se qualquer estudante for conduzido a um estudo sério de lógica matemática por este livro, ele terá servido ao principal propósito para o qual foi escrito.

Índice remissivo

A

[agregados](#)

alefe, [1-2](#), [3-4](#), [5-6](#), [7-8](#)

alguns, [1-2](#)

alio-relativa, *ver* anti-reflexiva

análise, [1-2](#)

ancestrais, [1-2](#), [3-4](#)

anti-reflexiva, [1-2](#)

argumento de uma função, [1-2](#), [3-4](#)

aritmetização da matemática, [1-2](#)

axioma multiplicativo, 118, [1-2](#)

[axiomas](#)

B

Bolzano, [1n](#)

botas e meias, [1-2](#)

C

campo de uma relação, [1-2](#), [3-4](#)

Cantor, Georg, 101, [1-2](#), [3n](#), [4](#), [5](#), [6](#), [7-8](#), [9-10](#)

casos, [1-2](#)

[classe mediana](#)

classe nula, [1](#), [2](#)

classes, [1](#), [2-3](#), [4-5](#)

lexivas, [1-2](#), [3-4](#), [5-6](#)

similares, [1-2](#)

[Clifford, W.K.](#)

coleções infinitas, [1-2](#)

conjunção, [1-2](#)

consecutividade, [1-2](#), [3-4](#)

constantes, [1-2](#)

construção, método de, [1-2](#)

contagem, [1-2](#)

continuidade, [1-2](#), [3-4](#)

cantoriana, [1-2](#)

dedekindiana, [1-2](#)

de funções, [1-2](#)
na filosofia, [1-2](#)
contradições, [1-2](#)
contrapartes objetivas, [1-2](#)
[convergência](#)
correlatores, [1-2](#)

D

Dedekind, [1](#), [2](#), [3n](#)
dedução, [1-2](#)
definição, [1-2](#)
[extensional e intensional](#)
derivadas, [1-2](#)
descrições, [1-2](#), [3-4](#), [5-6](#)
[dimensões](#)
disjunção, [1-2](#)
[diversidade](#)
domínio, [1-2](#), [3-4](#), [5-6](#)

E

entre, [1-2](#), [3-4](#)
equivalência, [1-2](#)
espaço, [1-2](#), [3-4](#), [5](#)
estrutura, [1-2](#)
[Euclides](#)
existência, [1](#), [2-3](#), [4-5](#)
exponenciação, [1-2](#), [3-4](#)
extensão de uma relação a, [1-2](#)

F

ficções lógicas, [1n](#), [2-3](#), [4-5](#)
[filosofia matemática](#)
[finito](#)
[fluxo](#)
forma, [1-2](#)
frações, [1](#), [2-3](#)
Frege, [1](#), [2n](#), [3](#), [4](#), [5n](#)
fronteira, [1](#), [2-3](#)
função de verdade, [1-2](#)
funções, [1-2](#)
 descritivas, [1-2](#), [3](#)
 intensionais e extensionais, [1-2](#)
 predicativas, [1-2](#)
 proposicionais, [1-2](#), [3](#), [4-5](#)

G

generalização, [1-2](#)
geometria, [1](#), [2-3](#), [4-5](#), [6-7](#), [8-9](#), [10-11](#)
analítica, [1-2](#), [3-4](#)

H

[Hegel](#)

I

idéias e proposições primitivas, [1](#), [2](#)
implicação, [1](#), [2](#)
[formal](#)
incomensuráveis, [1-2](#), [3](#)
incompatibilidade, [1-2](#), [3](#)
indiscerníveis, [1-2](#)
indivíduos, [1](#), [2](#), [3](#)
indução matemática, [1-2](#), [3](#), [4-5](#), [6-7](#)
inferência, [1-2](#)
infinitude, axioma da, [1n](#), [2](#), [3-4](#), [5-6](#)
[infinito](#)
cantoriano, [1-2](#)
de cardinais, [1-2](#)
de racionais, [1-2](#)
e séries e ordinais, [1-2](#)

[intervalos](#)

intuição, [1-2](#)
inverso, [1](#), [2-3](#), [4-5](#)
irracionais, [1-2](#), [3-4](#)
[irrealidade](#)

K

[Kant](#)

L

lacuna dedekindiana, [1-2](#), [3-4](#)
lei associativa, [1-2](#), [3-4](#)
lei comutativa, [1-2](#), [3-4](#)
lei distributiva, [1-2](#), [3-4](#)
Leibniz, [1](#), [2](#), [3](#)
Lewis, C.I., [1-2](#)
limite, [1](#), [2-3](#), [4-5](#)
de funções, [1-2](#)
lógica, [1](#), [2](#), [3-4](#)
matemática, [1-2](#), [3](#)
[logicização da matemática](#)

M

maior e menor, [1-2](#), [3-4](#)

mapas, [1](#), [2-3](#), [4](#)

matemática, [1-2](#)

máximo, [1](#), [2](#)

[Meinong](#)

método,

mínimo, [1](#), [2](#)

[modalidade](#)

multiplicação, [1-2](#)

N

[necessidade](#)

Nicod, [1-2](#), [3n](#)

nomes, [1](#), [2-3](#)

número cardinal, [1-2](#), [3-4](#), [5](#)

 complexo, [1-2](#)

 finito, [1-2](#)

 indutivo, [1](#), [2-3](#), [4-5](#)

 infinito, [1-2](#)

 irracional, [1](#), [2](#)

 máximo, [1-2](#)

 multiplicativo, [1-2](#)

 não-indutivo, [1](#), [2](#)

 natural, [1-2](#), [3](#)

 real, [1](#), [2](#), [3-4](#)

 lexivo, [1-2](#), [3-4](#)

 relação, [1](#), [2](#)

 serial, [1-2](#)

números de relação, [1-2](#)

números inteiros, positivos e negativos, [1-2](#)

O

O, [1](#), [2-3](#)

[Occam](#)

ocorrências, primárias e secundárias, [1-2](#)

ordem, [1-2](#)

[cíclica](#)

oscilação fundamental, [1-2](#)

P

[Parmênides](#)

particulares, [1-2](#), [3-4](#)

Peano, [1-2](#), [3-4](#), [5](#), [6](#), [7](#)

Peirce, [1n](#)

[permutações](#)

Pitágoras, [1](#), [2](#)

[Platão](#)
[pluralidade](#)
Poincaré, [1-2](#)
[pontos limitantes](#)
[pontos](#)
posteridade, [1-2](#)
 própria, [1-2](#)
postulados, [1-2](#)
 [precedentes](#)
 [premissas da aritmética](#)
 progressões, [1](#), [2-3](#)
 proposições, [1-2](#)
 analíticas, [1-2](#)
 elementares, [1-2](#)
 propriedades hereditárias, [1-2](#)
 propriedades indutivas, [1-2](#)
 [prova ontológica](#)

Q

quantidade, [1-2](#), [3](#)

R

razões, [1](#), [2-3](#), [4](#), [5](#)
reducibilidade, axioma da, [1-2](#)
 erente, [1-2](#)
 referido, [1-2](#)
relações assimétricas, [1](#), [2](#)
 conexas, [1-2](#)
 muitos-muitos, [1-2](#)
 muitos-um, [1-2](#)
 [quadrados de](#)
 lexivas, [1-2](#)
 seriais, [1-2](#)
 simétricas, [1-2](#), [3-4](#)
 similares, [1-2](#)
 transitivas, [1-2](#), [3-4](#)
 um-muitos, [1-2](#)
 um-um, [1-2](#), [3-4](#), [5-6](#)
representantes, [1-2](#)
[rigor](#)
[Royce](#)

S

seção dedekindiana, [1-2](#)
 fundamental, [1-2](#)
segmentos, [1](#), [2](#)

seleções, [1-2](#)

[semelhança](#)

[seqüentes](#)

séries, [1-2](#)

bem-ordenadas, [1](#), [2](#)

compacta, [1](#), [2](#), [3-4](#)

[condensada em si mesma](#)

dedekindiana, [1-2](#), [3-4](#)

fechada, [1-2](#)

[geração de](#)

infinitas, [1-2](#)

perfeitas, [1-2](#)

[Sheffer](#)

[silogismo](#)

símbolos incompletos, [1-2](#)

similaridade, de classes, [1-2](#)

de relações, [1-2](#), [3](#)

subclasses, [1-2](#)

[subtração](#)

sucessor de um número, [1-2](#), [3-4](#)

[sujeitos](#)

T

tautologia, [1](#), [2](#)

tempo, [1-2](#), [3](#), [4](#)

tipos lógicos, [1](#), [2-3](#), [4](#), [5](#)

todos, [1-2](#)

V

valor de uma função, [1](#), [2](#)

[valor de verdade](#)

variáveis, [1](#), [2](#), [3](#)

[Veblen](#)

[verbos](#)

[vizinhança](#)

W

Weiertrass, [1](#), [2](#)

[Wells, H.G.](#)

Whitehead, [1](#), [2](#), [3](#), [4](#)

Wittgenstein, [1n](#)

Z

Zermelo, [1-2](#), [3](#)

zero, [1-2](#)

Notas

Introdução

[1](#) Alusão ao poema “A Grammarian’s Funeral” de Robert Browning (1812-89). (N.T.)

A série dos números naturais

[1](#) *Principia Mathematica*, de Whitehead e Russell, Cambridge University Press, vol.I, 1910; vol.II, 1911; vol.III, 1913.

[2](#) Usaremos “número” nesse sentido no presente capítulo. Mais adiante a palavra será usada em sentido mais geral.

Definição de número

[1](#) A mesma resposta é dada de maneira mais completa e com maior desenvolvimento em Gottlob Frege, *Grundgesetze der Arithmetik*, vol.I, 1893.

[2](#) Como será explicado mais tarde, as classes podem ser consideradas ficções lógicas, fabricadas com base em características definidoras. Por enquanto, porém, simplificará nossa exposição tratá-las como se fossem reais.

[3](#) A noção de “relação um-um” corresponde à “função biunívoca”, também chamada “um a um”. As expressões “um muitos” e “muitos-um”, entretanto, não têm correspondência na nomenclatura atualmente utilizada. (N.R.T.)

Finitude e indução matemática

[1](#) Ver *Principia Mathematica*, vol.II. nota 110.

[2](#) Ver Capítulo 13.

[3](#) Para a geometria, na medida em que ela não é puramente analítica, ver *Principles of Mathematics*, parte VI; para a dinâmica racional, *ibid.*, parte VII.

[4](#) Essas definições e a teoria da indução generalizada se devem a Frege e foram publicadas já em 1879 em seu *Begriffsschrift*. Apesar do grande valor dessa obra, eu fui, acredito, a primeira pessoa a lê-la algum dia — mais de 20 anos após sua publicação.

A definição de ordem

[1](#) Relação “alio-relativa”, no original, Bertrand Russell usou esse termo por sugestão de C.S. Pierce. A nomenclatura atual é “anti-reflexiva”. (N.R.T.)

[2](#) Cf. *Rivista di Matematica*, IV, p.55ss; *Principles of Mathematics*, p.394 (§ 375).

[3](#) *Principles of Mathematics*, p.205 (§194), e referências dadas ali.

Similaridade das relações

[1](#) Isso não se aplica ao espaço elíptico, mas apenas a espaços em que a linha reta é uma série aberta. *Modern Mathematics*, editada por J.W.A. Young, p. 3-51 (monografia de O. Veblen sobre “The Foundations of Geometry”).

Números racionais, reais e complexos

[1](#) Vol.III, nota 300ss, especialmente 303.

[2](#) É claro que na prática continuaremos a falar de uma fração como (digamos) maior ou menor do que 1, querendo dizer maior ou menor do que a razão 1/1. Contanto que seja compreendido que a razão 1/1 e o número cardinal 1 são diferentes, não é necessário ter sempre o pedantismo de enfatizar a diferença.

[3](#) Estritamente falando, essa afirmação, bem como as que se seguem até o final do parágrafo, envolve o chamado “axioma da infinidade”, que será analisado num capítulo posterior.

[4](#) *Stetigkeit und irrationale Zahlen*, 2a ed., Brunswick, 1892.

[5](#) Para um tratamento mais completo do tema de segmentos e relações dedekindianas, ver *Principia Mathematica*, vol.II, nota 210. Para um tratamento mais completo de números reais, ver *ibid.*, vol.III nota 310ss, e *Principles of Mathematics*, caps. XXXIII e XXXIV.

Números cardinais infinitos

[1](#) Cf. *Principia Mathematica*, vol.II nota 123.

[2](#) Essa prova é tomada de Cantor, com algumas simplificações: ver *Jahresbericht der deutschen Mathematiker-Vereinigung*, I. (1892), p.77.

Limites e continuidade de funções

[1](#) *Principia Mathematica*, vol.II nota 230-234.

[2](#) Diz-se que um número é “numericamente menor” do que ε quando se situa entre $-\varepsilon$ e $+\varepsilon$ atualmente diz-se que o “valor absoluto do número é menor do que ε ” quando se situa entre $-\varepsilon$ e ε . [N.R.I.]

Seleções e o axioma multiplicativo

[1](#) Ver *Principia Mathematica*, vol.I nota 88. Também vol.III nota 257-8.

[2](#) *Mathematische Annalen*, vol. LIX, p. 514-6. Sob essa forma falaremos dele como o axioma de Zermelo.

O axioma da infinidade e os tipos lógicos

[1](#) Sobre esse assunto, ver *Principia Mathematica*, vol.II p.120. Sobre os problemas correspondentes no tocante a razões, ver *ibid.*, vol.III p.303.

- [2](#) Vol.I, Introduction, cap. II nota 12 e nota 20; vol.II, Prefatory Statement.
- [3](#) “Mathematical Logic as based on the Theory of Types”, vol. XXX, 1908, p. 222-62.
- [4](#) “Les paradoxes de la logique”, 1906, p. 627-650.
- [5](#) Bolzano, *Paradoxien des Unendlichen*, 13.
- [6](#) Dedekind, *Was sind und was sollen die Zahlen?*, n. 66.

Incompatibilidade e a teoria da dedução

- [1](#) Usaremos as letras p, q, r, s, t para denotar proposições variáveis.
- [2](#) O termo foi cunhado por Frege.
- [3](#) *Trans. Am. Math. Soc.*, vol.XIV, p. 481-8.
- [4](#) *Proc. Camb. Phil. Soc.*, vol.XIX, I, janeiro de 1917.
- [5](#) Nenhum princípio desse tipo é enunciado em *Principia Mathematica* ou no artigo de M. Nicod mencionado. Mas isso parece ser uma omissão.
- [6](#) Ver *Mind*, vol.XXI, 1912, p. 522-31; e vol.XXIII, 1914, p. 240-47.

Funções proposicionais

- [1](#) O método de dedução é dado em *Principia Mathematica*, vol.I p. 9.
- [2](#) Por razões lingüísticas, para evitar sugerir o plural ou o singular, muitas vezes é mais conveniente dizer “ ϕx não é sempre falso” do que “ ζ às vezes” ou “ ζ é verdadeiro às vezes”.

Descrições

- [1](#) *Untersuchungen zur Gegenstandstheorie und Psychologie*, 1904.

Classes

- [1](#) Ver *Principia Mathematica*, vol.I, p.75-84 e 20.
- [2](#) O leitor que deseje uma discussão mais completa deveria consultar *Principia Mathematica*, Introduction, cap. II; também * 12.

Matemática e lógica

- [1](#) As proposições primitivas em *Principia Mathematica* são tais que permitem a inferência de que existe pelo menos um indivíduo. Mas vejo isso agora como uma imperfeição na pureza lógica.
- [2](#) A importância de “tautologia” para uma definição da matemática me foi mostrada por meu ex-aluno Ludwig Wittgenstein, que estava trabalhando com o problema. Não sei se ele o resolveu, e nem sequer se está vivo ou morto.

Título original:

Introduction to Mathematical Philosophy

Tradução autorizada da edição inglesa publicada por Taylor & Francis Books, de Londres, Inglaterra. Esta obra foi originalmente publicada por George Allen & Unwin.

Copyright © 1996, Bertrand Russell Peace Foundation

Copyright da Introdução © 1993, John G. Slater

Copyright da edição brasileira © 2007:

Jorge Zahar Editor Ltda.

rua Marquês de São Vicente 99 - 1º andar

22451-041 Rio de Janeiro, RJ

tel.: (21) 2529-4750 / fax: (21) 2529-4787

editora@zahar.com.br

www.zahar.com.br

Todos os direitos reservados. A reprodução não-autorizada desta publicação, no todo ou em parte, constitui violação de direitos autorais. (Lei 9.610/98)

Capa: Miriam Lerner

Edição digital: abril 2011

ISBN: 9788537804131

Arquivo ePub produzido pela [Simplíssimo Livros](#) - [Simplicissimus Book Farm](#)
